

**‘인공지능(AI)
윤리 가이드라인’의
중요성과 국가별
대응 현황
: 국내**



'인공지능(AI) 윤리 가이드라인'의 중요성과 국가별 대응 현황 : 국내

1. 인공지능 윤리(AI Ethics)

■ 인공지능 윤리의 개념

○ '인공지능 기술 전망과 혁신정책 방향' (과학기술정책연구원, 정책연구 2019-13)에서는 인공지능 윤리(AI Ethics)를 “인공지능 관련 이해관계자들이 준수해야 할 보편적 사회 규범 및 관련 기술”로 조작적 정의하고 있다.

■ 인공지능 윤리 논의의 배경

인간을 위협하는 인공지능의 출현은 더 이상 영화 속 상상이 아니다. 현실에서 벌어지고 있는 일이다. 탐사보도 전문매체 프로 퍼블리카는 범죄 전과자의 얼굴 이미지를 기반으로 재범률을 예측하는 인공지능 알고리즘을 테스트했다. 결과적으로 인공지능은 흑인의 재범률을 백인에 비해 실제보다 훨씬 더 높게 추론했다. 인공지능이 판사를 대체할 경우, 인간은 기계에 의해 인종차별을 당할 수 있는 것이다.

미국 스탠포드 대학이 2017년 발표한 연구도 주목할 만하다. 인공지능에 온라인 데이트 사이트에 공개된 남녀 프로필 사진으로 성적 지향성을 추론하게 한 결과, 인공지능은 사람보다 더 정확한 판단을 내렸다. 성소수자가 감당해야 하는 사회적 불이익을 감안하면, 인공지능이 내놓은 값의 신뢰성을 따지기 전에 이러한 연구 자체도 비윤리적일 수 있다.

인공지능의 학습 알고리즘은 방대한 데이터를 기초로 하는 일종의 통계적 추론이다. 여성, 장애인, 유색인종 등 소수 집단의 데이터는 수집이 어렵기 때문에, 알고리즘이 학습하는 데이터에도 편향성이 생길 수 있다. 기업 인사팀이 수많은 구직자들의 이력서를 1차적으로 필터링하기 위해 인공지능을 도입한다면, 구직자들은 인공지능에 의한 차별을 당하게 되는 것이다.

■ 인공지능 윤리 침해 사례

- 인공지능을 비윤리적으로 활용하는 대표적인 경우는 대량의 인명 피해가 예상되는 자율살상 무기 개발
 - 2025년 기준으로 인공지능을 탑재한 군사용 로봇과 드론 시장 규모는 민간 시장 규모를 압도할 것으로 예상
 - 2019년 1월 예멘의 한 공군기지에서 개최된 정부군 행사에 후티(Houthi) 반군의 드론이 폭파해 정부군 6명이 사망하

고 관료 12명이 부상하는 사고 발생

○ 의도적으로 가짜 이미지나 영상, 뉴스, 음성을 생성해 배포하는 행위도 만연

- 딥페이크라고 불리는 이러한 편집물들은 개인의 평판을 훼손시키고 사회적 구성원 간의 신뢰를 저하
- 이스라엘의 벤구리온(Ben-Gurion) 대학교 연구진은 딥러닝을 이용해 의료 영상을 조작한 실험 결과를 2019년 4월 논문으로 발표

2. 국내 기업의 AI 윤리 대응 동향

■ 인공지능에 대한 기업의 사회적 책임

이러한 인공지능의 폐해를 막기 위해서는 기업의 윤리의식이 선행되어야 한다. 인공지능의 판단은 알고리즘에 기반하며, 알고리즘은 기업에 속한 인간 개발자가 설계하기 때문이다. 이를 위해서는 윤리규정과 이를 준수하려는 조직의 의지가 필수다.

■ 카카오

○ 2018년 1월 31일, 국내 업계 최초로 인공지능 알고리즘 윤리 헌장을 발표했다.

| 카카오의 윤리 헌장 |

원칙	내용
1. 카카오 알고리즘의 기본원칙	카카오는 알고리즘과 관련된 모든 노력을 우리 사회 윤리 안에서 다하며, 이를 통해 인류의 편익과 행복을 추구한다.
2. 차별에 대한 경계	알고리즘 결과에서 의도적인 사회적 차별이 일어나지 않도록 경계한다.
3. 학습 데이터 운영	알고리즘에 입력되는 학습 데이터를 사회 윤리에 근거하여 수집·분석·활용한다.
4. 알고리즘의 독립성	알고리즘이 누군가에 의해 자의적으로 훼손되거나 영향받는 일이 없도록 엄정하게 관리한다.
5. 알고리즘에 대한 설명	이용자와의 신뢰 관계를 위해 기업 경쟁력을 훼손하지 않는 범위 내에서 알고리즘에 대해 성실하게 설명한다.

카카오의 윤리헌장은 최근 국내외에서 활발히 논의되고 있는 사회적 차별 금지, 데이터 윤리, 알고리즘의 독립성 및 설명 의무 등에 대한 원칙을 포함하고 있다. 즉 알고리즘 기술이 인간의 도덕적 가치와 윤리 관점에 부합되어야 하고 궁극적으로는 인간의 삶의 효용성을 높일 수 있어야 한다는 점을 강조하고 있다.

■ 삼성전자

○ 2018년 국내기업으로는 최초로 인공지능 국제 컨소시엄인 PAI(Partners on AI)에 가입

- 인공지능 미래에 대한 범사회적인 논의를 진행 중인 삼성전자는 향후

AI 안전성(Safety-Critical AI)과 AI 공정성·투명성·책임성(Fair, Transparent, and Accountable AI), AI의 사회적 영향(Social and Societal Influences of AI) 등 다양한 분야에 참여를 확대할 계획이다.

구체적으로는 각각의 분야에서 의료·안전 등의 영역에서 윤리적인 의사결정 방법, AI 기술 개발 시 고려할 윤리적 요소, 자동화로 인한 직업 소멸 등으로 인한 분쟁 방지, 인간의 인지능력 향상을 위한 AI 정보 제공, 사생활, 민주주의, 인권 등 AI가 사회에 미치는 영향, 교육, 주택, 공중 보건 등 공공재로서의 AI 등을 연구한다.

■ 네이버

- 인공지능 기반 서비스에서 안전한 데이터 활용을 위해 'Privacy By Design 원칙'을 적용
 - 서비스 출시 단계에서부터 높은 프라이버시 보호 기준으로 데이터 활용에 있어 법령 위반사항이나 이용자 프라이버시 보호 관점에서 문제가 없는지를 검토하고 있다.
- '네이버 프라이버시센터' 홈페이지, 개인정보 보호 블로그 및 SNS를 통해 이용자 프라이버시 보호를 위한 활동과 관련 정보를 공개하고 이용자의 개선제안 등을 받고 있다.

'네이버 프라이버시 센터' 홈페이지에서는 네이버의 개인정보 보호원칙, 서비스 운영 관련 정책, 보호 활동, 투명성 보고서, 개인정보 보호와 활용 논의를 위한 Privacy White Paper 등 프라이버시 관련 보고서를 제공하고 있다

| 네이버의 프라이버시 관련 보고서 |

구분	내용	발행주기
투명성 보고서 통계	네이버(주)에서 수사기관에 제공한 이용자 정보 통계	연 2회
개인정보보호 리포트	네이버(주)의 개인정보 보호 활동	연 1회
Privacy White Paper	'이용자 프라이버시 보호' 관련 전문 연구 수행 결과	연 1회

자료: 네이버 프라이버시센터 홈페이지

■ 한국인공지능협회

- 2012년 스타트업 실무자들의 커뮤니티로 시작하여 2017년 사단법인으로 설립. "산업의 지능화"와 "AI 기술의 대중화"를 목표로 인공지능 기술 기업 및 유관기관들을 중심으로 네트워크를 구축
- 인공 지능 산업에 필요한 윤리의식과 교육의 필요성을 논의하기 위한 목적으로 2019년 7 월 '제1회 인공지능 윤리+교육 포럼'을 개최하였으며, 포럼 이후 메니페스토의 일환으로 협회 클러스터 기업 72개 업체가 참여한 2019 인공지능 윤리+교육 포럼 공동선언문을 채택했다.

| 2019 인공지능 윤리+교육 포럼 공동선언문 |

본 선언문은 대한민국 인공지능 기술의 지속적 성장과 안전한 발전을 위한 윤리적 연구 개발 뿐 아니라, 이와 연계하여 올바른 인공지능 교육을 할 수 있도록 개인과 기업, 국가 그리고 세계기구에 실천적 가이드라인을 제시하는 데 그 목적이 있다.

1. 인공지능 기술은 처음부터 '사람 우선' 원칙으로 인류에게 육체적·정신적·환경적으로 도움을 줄 수 있도록 연구·개발되어야 하고 동시에 설명될 수 있어야 한다. 또한, 우리가 그것을 신뢰하고 공정하게 공유하여 인류 공동의 선을 실현할 수 있도록 해야 한다.
2. 국민 개인은 '자기의 기술(The Technology of the Self)'을 학습하여 개인 데이터를 스스로 관리해야 한다. 또 미래 인공지능 기술이 가져올 직업적 변화에 대비하여 자신의 적성과 능력을 파악, 소비자인 동시에 생산자가 될 수 있는 새로운 디지털 기술을 습득하고 그것을 발전시켜야 한다.

3. 모든 기업은 시작단계에서부터 '안전 우선'을 원칙으로 관계자 교육훈련을 포함한 인공지능 개발과정을 설계하고, 해당 제품의 사람과 사회에 대한 영향을 종합적으로 분석하여 차별과 편견이 없도록 제작해야 한다. 그리고 비상시를 대비하여 위험요소를 즉각 제거할 수 있는 통제기술을 해당 제품에 포함해야 하며, 제품의 전체 제작과정을 표준화시키고, 사후 관리에 책임을 지며, 영업비밀을 훼손하지 않는 범위에서 그것을 설명할 수 있어야 한다.
4. 국가 교육기관은 국민 각자가 국민공동기본교육과정을 통해 생애주기별로 필요한 '스마트 라이프(Smart Life)' 활용기술을 습득하여 일상 생활능력을 확장할 수 있게 하며, 직업적으로 개인의 역량을 최대로 향상시킬 수 있는 '자기 계량화(Quantified Self)' 기술을 포함한 '인간공학(Human Engineering)'에 대한 학습기회를 제공해야 한다.
5. 국가는 국민의 인권과 재산, 프라이버시를 보호하기 위해 산업별 인공지능 기술의 긍정적·부정적 영향을 전체적으로 조사·분석·예측하여 그 결과를 공유하고, 관계자에 대한 자격부여 및 교육·훈련을 해야 한다. 특히, 종합적 제품관리를 위해 기업 경쟁력을 침해하지 않는 범위에서 산업별 호준화된 평가 매트릭스를 적용하여 투명성과 타당성을 확보하고, 합리적 개입과 감독을 포함한 국가 차원의 가버넌스시스템을 구축해야 한다.
6. 인공지능 기술을 그 경계가 없이 세계적으로 확장될 수 있으므로 개인과 기업은 세계시민으로서의 역할교육을 받아야 하며, 한 나라의 국가적 가버넌스 시스템은 UN과 같은 세계기구의 글로벌 시스템과 연동시켜 지속해서 발전해 나가야 한다.
7. UN등 세계기구는 인공지능 기술이 인간의 생리적 진화를 앞지를 수 있는 문명사적 변혁에 대비하여, 인간의 존엄성과 그 존재론적 의미와 가치를 재발견해야 한다. 나아가서는 '사람 우선'의 원칙을 기반으로 인공지능과 인간의 표준화된 기술적 협업 방법을 개발하고 보급하여 상호갈등을 없애고, 세계평화와 질서를 유지할 수 있는 정책을 강구해야 한다.

자료: 한국인공지능협회 홈페이지

3. 국내 학계의 AI 윤리 대응 동향

■ 고등과학원 초학제연구단

○ 고등과학원은 순수이론기초과학 연구기관으로서 학제간 교류와 소통을 통해 기존 연구방법에 구애받지 않고 새로운 연구주제와 방법을 생성하는 초학제 연구프로그램을 2012년부터 진행하고 있다.

- 2017년부터는 초학제 연구프로그램 주제로 인공 지능을 정하고, 2017년부터 2018년까지는 <인공지능: 과학, 역사, 철학>, 2019년부터는 <인공지능과 포스트휴머니즘> 제목으로 인공지능 현장 연구자와 인문학자들이 모여 인공지능과 사회적 영향에 대한 발표와 토론을 진행하는 워크숍을 진행하고 있다.

2018년 10월 18일에 진행된 <로봇윤리의 철학적 기초> 워크숍에서 서울 시립대 이종원 교수는 인간 친화적인 형태로 개발되는 AI 로봇 등을 예로 들며 자율성을 가진 AI 로봇에 대한 논의가 필요하며, 책임 귀속의 최종 주체는 인간이 될 수 밖에 없지만, 사람들의 집합체가 만들어낸 인공물이 인간에게 영향을 준다면 사회적 차원의 합의가 필요하다고 설명하였다. 또한, 윤리는 기술개발 과정에서 자율적인 자기 규제를 가능하게 하므로 법이 못하는 부정적(negative) 규제 허용 범위를 제시하는 역할을 할 수 있을 것이라고 전망하였다.

■ 한국인공지능법학회

○ 2019년 3월부터는 인공지능 기술과 법제도, 윤리, 그리고 정책의 상호작용에 관한 문제를 종합적 시각에서 풀어나가고 관련 정책에 대한 사회적 공론의 장을 마련할 취지로서, AI 분야별 정·관·산·학·연 전문가 및 언론사, 시민단체 등 50명 내외로 구성된 AI 정책포럼을 발족

○ 2019년 2월에는 ‘인공지능과 법’이라는 책을 출간하였으며, 이 책은 인공지능 관련법 패러다임 변화, 개별법, 정보보호법, 경쟁법, 사회보장법 등 법적 이슈와 자율주행차, 의료, 킬러로봇, 로보어드바이저 등 기술분야에서의 윤리적·법적 쟁점 등을 분석하고 이에 대한 대안을 제시하고자 하였다.

■ 한국인터넷윤리학회

○ 2017년에 ‘인공지능 및 로봇기술의 윤리적 쟁점’ 특별 세미나 등을 개최하는 등 인공지능 분야의 윤리 이슈에 대한 학술 활동을 진행

○ 2017년 논문지 ‘The Digital Ethics’를 발간하여 인공지능 윤리에 대한 다수의 논문을 게재하고 있으며, 정부 부처와 AI 윤리 관련 정책연구도 진행하고 있다.

- 2017년에는 지능정보사회 발전을 위한 윤리 원칙 관련 위탁연구 과제(방송통신위원회)를 수행하였고, 2018년에는 전자정부에 인공지능 기술 활용 시 윤리기준 관련 연구(행정안전부)를 수행하였다.

■ 정보통신정책연구원

○ 2018년 10월 발간한 “4차 산업혁명시대 산업별 인공지능 윤리의 이슈분석 및 정책 적 대응방안 연구”는 지능형로봇윤리현장, 지능정보사회윤리가이드라인, 지능정보사회윤리현장 등 일반적/총론적 규범 차원의 논의에 그치고 있는 인공지능 윤리를 인공지능 기술 및 서비스의 구체적 적용 분야의 특성에 따른 윤리이슈 검토 및 대응방안 논의로 구체화해야 할 필요성을 지적하였다.

- 본 연구에서는 AI 기술 적용 및 윤리이슈 쟁점 정도를 고려하여 자율주행자동차(제조), AI의료서비스(의료), 로보 어드바이저(금융), 킬러로봇/치명적 자율무기체계(LAWS: Lethal Automotive Weapons System)(국방(군사))의 윤리이슈 분석 및 대응방안을 도출하였다.

4. 정부의 AI 윤리 대응 동향

■ 방송통신위원회

○ 2019년 11월 11일 정보통신정책연구원과 함께 AI 시대 정부, 기업, 이용자 등 구성원이 지켜야 할 '이용자 중심의 지능정보서비스 기본 원칙' 발표.

- 이는 AI 시대 이용자의 권리와 이익이 충분히 보호될 수 있도록 정부·기업·이용자 등 구성원들이 함께 지켜가야 할 기본적인 원칙으로, 주요 기업과 분야별 전문가들이 자문에 참여했다. 2018년에는 지능정보화 이용자 포럼을 통해 원칙 마련을 위한 기초 연구가 이루어졌으며, 2019년 10월까지 원칙의 주요내용에 대한 ICT 기업*과 전문가 의견수렴 과정을 거쳐 완성되었다.

(* 구글코리아, 페이스북코리아, 넷플릭스, 카카오, 삼성전자, KT, SKT, LGU+, 한국 IBM, 한국 마이크로소프트, 솔트룩스, 인텔코리아, BSA Korea 등)

○ 지능정보서비스 기본 원칙 주요내용

- (사람 중심의 서비스 제공) 지능정보서비스의 제공과 이용은 사람을 중심으로 그 기본적 자유와 권리를 보장하고 인간의 존엄성을 보호할 수 있는 방향으로 이루어져야 한다.

- (투명성과 설명가능성) 지능정보서비스가 이용자에게 중대한 영향을 끼칠 경우 기업의 정당한 이익을 침해하지 않는 범위에서 이용자가 이해할 수 있도록 관련 정보를 작성해야 하며, 이용자 기본권에 피해를 유발했을 때 예측, 추천, 결정의 기초로 사용한 주요요인을 설명할 수 있어야 한다.
- (책임성) 지능정보사회의 구성원들은 지능정보서비스의 올바른 기능과 사람 중심 가치의 보장을 위한 공동의 책임을 인식하고, 관련한 법령과 계약을 준수한다.
- (안전성) 안전하고 신뢰 가능한 지능정보서비스의 개발과 이용을 위해 모두가 노력하고, 지능정보서비스가 초래할 수 있는 피해에 대한 자율적인 대비체계를 서비스 제공자와 이용자가 수립하고 운영한다.
- (차별금지) 지능정보서비스가 사회적·경제적 불공평이나 격차를 초래할 수 있다는 점을 인식하고, 알고리즘 개발과 사용의 모든 단계에서 차별적 요소를 최소화할 수 있도록 노력한다.
- (참여) 지능정보사회의 구성원들은 공적인 이용자 정책 과정에 차별 없이 참여할 수 있으며, 공적 주체는 제공자와 이용자가 실질적으로 의견을 제시할 수 있는 정기적인 통로를 조성해야 한다.
- (프라이버시와 데이터거버넌스) 지능정보서비스의 개발, 공급 및 이용의 전 과정에서 개인정보 및 프라이버시를 보호 하며, 구성원들은 기술적 이익의 향유와 프라이버시 보호 사이의 균형을 위해 지속적인 의견 교환에 참여한다.
- 이용자 보호를 위한 공동의 노력
 - 지능정보사회의 구성원들은 공동의 기본 원칙에 입각하여 지능정보사회의 기본 가치를 수호할 수 있도록 자율적인 노력을 지속한다.
 - 지능정보사회 이용자 보호를 위한 실질적이고 지속적인 논의를 위해 지능정보사회의 구성원들이 참여하는 협의회를 구성하여 운영한다.

■ 정부

○ 2019년 12월 17일 정부는 문재인 대통령 주재로 열린 제 53회 국무회의에서, 과학기술정보통신부를 비롯한 전 부처가 참여하여 마련한「인공지능(AI) 국가전략」을 발표

- ▶ AI 반도체 세계 1위, 전국 단위 AI 거점화 등 세계를 선도하는 AI 생태계 조성
- ▶ 전 생애·모든 직군에 걸친 AI교육 실시 및 세계 최고의 AI인재 양성
- ▶ 현 전자정부를 차세대 지능형 정부로 대전환하여 국민 체감도 향상
- ▶ 사회보험 확대 등 일자리 안전망 확충 및 AI 윤리 정립으로 사람 중심 AI 실현
- AI를 통해 경제효과 최대455조원 창출, 삶의 질 세계 10위 도약(~30) -

| AI 윤리 규범 및 가이드라인 마련 현황 |



○ 추진 전략

- AI 등 지능형 기술을 기반으로 사이버 위협 대응시스템을 고도화하고, 역기능 대응을 위한 기술개발 및 범부처 협업 체계 구성
- 사람 중심의 AI, 인간다운 AI 구현을 위해 사회적 공론화와 공감대 형성을 바탕으로 글로벌 수준의 AI 윤리규범을 정립

○ 역기능 방지 및 AI 윤리체계 마련

- ① AI 기반 사이버침해 대응체계 고도화('20~)
- ② 딥페이크* 등 신유형의 역기능 대응을 위한 범부처 협업체계 구축('20)
(* AI 기반 영상 합성기술 또는 그 영상, 신시장 창출과 동시에 명예훼손 등 부작용도 우려)
- ③ AI 신뢰성·안전성 등을 검증하는 품질관리체계 구축 추진('20~)
- ④ OECD 등 글로벌 규범에 부합하는 AI 윤리기준 확립('20) 및 AI 윤리교육 커리큘럼* 개발·보급('21~)
(* (학생·이용자) AI와 생명윤리, 개인정보보호 / (개발자) 윤리적 AI 설계, 정보보안 등)
- ⑤ 이용자 보호를 위한 중장기적 정책 수립 지원체계 마련

○ 정부는 대통령 직속의 현 4차산업혁명위원회를 AI의 범국가 위원회로 역할을 재정립하여 이번 전략의 충실한 이행을 위한 범정부 협업체계를 구축해 나갈 방침

- 특히, 대통령 주재 전략회의를 개최하여 전 국민 교육, 전 산업 AI 활용 등 범정부적 과제의 실행력을 확보하고, 대국민 성과 보고대회도 병행하여 국민의 참여와 성과 확산에도 노력할 계획

■ 금융위원회

○ 금융 분야 인공지능(AI) 윤리 가이드라인 및 인프라를 조성해 AI 기술 활성화에 나설 계획.

- 금융위원회는 2020년 7월 16일 국내·외 금융 분야 AI 활용 및 정책 동향 파악을 목표로 '금융 분야 인공지능(AI) 활성화' 워킹 그룹을 조직, 첫 회의를 개최

금융 분야 AI 활성화 정책을 추진할 수 있도록 워킹 그룹을 구성하고, 약 4개월간 운영해 '금융 분야 AI 활성화 방안'을 마련할 계획이다.

워킹 그룹은 ▲AI 관련 규제 ▲AI 인프라 ▲소비자 보호 ▲레그테크*와 쉐테크* 접목 4개 측면에서 AI 활성화 정책을 추진한다.

○ AI 활성화를 목표로 관련 규제 개선 및 규율 체계를 정립.

- 가이드라인 형태로 AI 금융 서비스 개발에 특화된 실무 프로세스를 마련하고, 과학기술정보통신부 AI 법제정비단과 협력해 적법성·공정성을 담은 '금융 분야 AI 윤리 가이드라인'을 갖춘다.

○ AI 활용 시 소비자 권한을 보호할 수 있는 체계를 확보.

- AI 업무 처리로 발생하는 소비자 피해를 해결할 수 있도록 책임 주체와 구제 절차 등 기준을 마련한다. AI가 도출한 결과를 객관적으로 설명할 수 있는 '설명가능 AI(XAI)' 기준도 정립한다.

(*레그테크는 '레귤레이션(규제)'과 기술을 의미하는 '테크놀로지(기술)'의 합성어로 AI를 활용해 복잡한 금융 규제를 기업이 쉽게 이해하고 지킬 수 있도록 하는 기술이다. *쉐테크는 '슈퍼비전(감독)'과 '테크놀로지(기술)'의 합성어로 신 기술을 활용해 금융 감독 업무를 효율적으로 수행하는 기술이다.)

■ 과학기술정보통신부

○ 2020년 9월 2일 디지털 뉴딜의 핵심, '데이터 댐'사업 본격 착수

지난 7월 발표된 『디지털 뉴딜』 대표과제인 '데이터 댐' 프로젝트의 7대 핵심사업들을 수행할 주요기업 등의 선정 작업을 마무리하고, 본격적으로 사업을 추진한다.

'데이터 댐' 7대 핵심사업은 미국 대공황 시기의 '후버댐' 건설과 같은 일자리와 경기부양 효과에 더하여 우리 미래를 위한 투자와 각 분야의 혁신을 동시에 추진하기 위해 기획된 ①인공지능 학습용 데이터 구축, ②인공지능 바우처와 ③인공지능 데이터 가공바우처 사업, ④인공지능 융합 프로젝트(AI+X), ⑤클라우드 플래그십 프로젝트, ⑥클라우드 이용바우처 사업, ⑦빅데이터 플랫폼 및 센터 구축 등 7개 사업이다.

○ AI 법제도 개선 및 윤리 정립

향후 계획으로는 AI로 인한 경제·사회 전반의 패러다임 변화에 대응하고자 규제개선 사항을 종합하여 AI분야 법제도 개선 로드맵을 제시할 계획이다.(20.11월 잠정) 아울러 금년 12월 AI시대의 기본법제인 『지능정보화 기본법』이 시행됨에 따라 하위법령을 완비하고, 주요국, 국제기구 등의 AI 윤리규범을 비교 분석하여 우리나라의 AI 윤리 기준도 정립할 예정이다.

| 향후계획(후속 조치사항) |



출처

‘인공지능 기술 전망과 혁신정책 방향’ (과학기술정책연구원 /이제영,김단비,양희태/ 정책연구 2019-13)

<http://www.stepi.re.kr/app/report/view.jsp?cmsCd=CM0012&categCd=A0201&ntNo=933>

방통위, 「이용자 중심의 지능정보사회를 위한 원칙」 발표

<https://kcc.go.kr/user.do?mode=view&page=A050300000&dc=K00000200&boardId=1113&boardSeq=47874>

금융위 「금융분야 인공지능(AI) 활성화」 계획

<http://m.fsc.go.kr/info/menu1/detail.do?oneDepth=1&twoDepth=1&tabDepth=&idx=34046&focus=0&pageIndex=1&searchStartDate=&searchEndDate=&searchOption1=1&searchOption=>

과기정통부 「인공지능(AI) 국가전략」 발표

https://www.msit.go.kr/web/msipContents/contentsView.do?catelId=_policycom2&artId=2405727

과기정통부 「디지털 뉴딜의 핵심, 데이터 댐 사업 본격 착수」

https://www.msit.go.kr/web/msipContents/contentsView.do?catelId=_policycom2&artId=3070219