



보도시점 2024. 11. 26.(화) 10:00
(2024. 11. 26.(화) 석간)

배포 2024. 11. 25.(월) 16:00

과기정통부, '인공지능 신뢰·안전성 학술회의' 개최, 책임있는 인공지능 위한 민·관 성과와 이행현황 공유

- 세계적 인공지능 석학 요슈아 벤지오 교수, 인공지능 안전 확보를 위한 정책과 정책 추진 전략 기조강연
- '인공지능 |서울기업서약('24.5.) 참여 국내 6개사, 인공지능 안전 확보 자발적 이행현황 발표
- '생성형 인공지능 잠재적 위험성 확인 경진대회(레드팀 챌린지)('24.4.) 분석결과 소개, 인공지능 안전 확보를 위한 업계활용 지원
- '제2회 인공지능 신뢰성 대상 선정 다비오(Eartheye 2.0), 액션파워(daglo) 등 총 8점 수상

과학기술정보통신부(장관 유상임, 이하 '과기정통부')는 11.26일(화) 피스앤피크 컨벤션(서울 전쟁기념관)에서 국내를 대표하는 인공지능 분야 대기업·신생기업, 연구자 등 200여명이 참여하는 가운데, '인공지능 신뢰·안전성 학술회의'를 개최하였다.

< '인공지능 신뢰·안전성 학술회의' 행사개요 >

- (주최/주관) 과기정통부 / 한국정보통신기술협회(TTA), 정보통신정책연구원(KISDI), 전자신문
- (기조강연) 퀘백 인공지능연구소 요슈아 벤지오 교수(국가인공지능위원회 국제 자문 위원), 카이스트 오혜연 인공지능 연구원장(국가인공지능위원회 안전·신뢰분과장)
- (주요참석기관) 한국정보통신기술협회, 네이버 카카오, SKT, KT, 삼성전자 LG AI연구원 포티튜마루, 튜닉 셀렉트스타

그간 과기정통부는 '국가 인공지능 윤리기준'('20.12.)을 토대로, 제3의 전문기관을 통한 민간자율 '인공지능 신뢰성 인증'('23.11.~) 지원, '생성형 인공지능 잠재적 위험성 취약점 발굴 경진대회(AI레드팀 챌린지)' 개최('24.4.), 국제 인공지능 규범가치(안전혁신포용)에 대한 합의를 이끌어낸 '인공지능 서울정상회의'('24.5.)의 성공적 개최, '국제 인공지능 안전연구소 협력망' 참여('24.11.20~21.), 인공지능 안전연구소 출범('24.11.27. 예정) 등 지속가능한 인공지능 혁신생태계 조성을 뒷받침하기 위해 인공지능 신뢰·안전성에 대한 정책적 지원을 강화하고 있다.

이번 학술회의에서는 최근 국제적 차원에서 강조되고 있는 인공지능 신뢰·안전성 관련 기술·정책 동향을 공유하고, '24년 정부지원 연구결과물의 주요 성과와 민간의 이행현황을 점검하며, 국내 인공지능 신뢰·안전성 확보의 우수사례를 선정·확산하기 위한 목적으로 개최되었다.

이번 학술회의 행사는 크게 4가지 주제로 진행되었다.

① 국내·외 인공지능 신뢰·안전 기술·정책 동향 공유

인공지능 분야 세계적 석학인 몬트리올대 요슈아벤지오 교수는 기조강연에서 “인공지능안전 확보를 위한 정책과 과학의 중요성을 강조하며, 최첨단 인공지능 모형의 위험에 대한 효과적인 관리를 위해 국내법-국제협약간의 상호조화가 필요하다”고 하였다. 또한, “과학측면에서는 인공지능 모형의 정렬과 통제가 중요쟁점으로, 정량적으로 측정가능한 인공지능 모형의 위험평가, 위험관리기술 확보를 위한 정부지원 확대와 인공지능 안전연구소의 역할이 필요하다”고 강조하였다.

이어진 두 번째 기조강연에서 한국과학기술원(KAIST) 오혜연 인공지능 연구원장은 국제사회의 인공지능 패권 경쟁 동향을 공유하고, 전략자산으로서 인공지능의 중요성을 강조하였다. 아울러, 인공지능 기술의 언어·문화적 포용성, 격차 문제 등 글로벌 사회에 미치는 영향을 소개하고, 인공지능 안전분야 신시장 창출과 기업 경쟁력 확보관점에서 인공지능 안전 정책의 추진전략을 소개하였다.

최근, 국내업계에서도 인공지능 신뢰·안전성 확보를 위한 기술개발과 평가 기법, 데이터셋 구축 등 다양한 연구가 진행되는 가운데, 튜넵 박규병 대표는 ‘인공지능 안전장치(가드레일) 해결책을 활용한 다양한 형태의 위험 탐지 및 대응방안’을 소개하였고, 포티투마루 김동환 대표는 ‘인공지능 신뢰성 인증(한국정보통신기술협회) 획득 등을 통한 자사 생성형인공지능 모형(LLM) 신뢰·안전성 확보사례’를 주제로 발표하였다.

② '인공지능 서울기업서약'(24.5.) 이행현황 발표(국내 6개사)

지난 5월에 개최된 '인공지능 국제 토론회(AI글로벌포럼)'에서 국내·외 14개 기업*은 '인공지능 서울기업서약'에 참여하여 책임있는 인공지능 개발·활용 보장 등에 대한 자발적 이행을 약속**한 바 있다.

* (서약참여기업) [국내] 삼성전자, NAVER, 카카오, SKT, KT, LG AI연구원
[국외] Google, OpenAI, Anthropic, IBM, Salesforce, Cohere, MS, Adobe

** (서약주요내용) "우리는 (i)취약성과 위험을 식별·평가하고, (ii)내부 평가결과를 인공지능 모형 수명 주기에 통합하며, 적절한 경우 인공지능 안전연구소 같은 국내 정부기관의 평가와 같은 외부 평가 결과를 고려하고, (iii)인공지능 안전사고를 점검하고 대응하는 강력한 내부 협치 및 위험관리 정책을 마련할 것을 약속"

이날 학술회의에서는 '인공지능 서울기업서약'에 참여한 국내 6개 기업이 모두 참여한 가운데, 안전하고 신뢰할 수 있는 인공지능을 개발하기 위한 위험 관리방안 수립, 기술 연구, 내부 협치체계 마련 등 각사의 이행현황*을 공유하였다.

* (네이버) '인공지능 안전성 체계(AI Safety Framework, ASF)' 수립('24.6.)
(카카오) 인공지능 위험관리체계(Kakao AI Safety Initiative) 구축·운영('24.10.~)
(삼성전자) 내장형 인공지능에 대한 신뢰·안전성 수립 및 운영('24.10.)
(SKT) 인공지능 협치체계(AI 거버넌스) 기본원칙을 구체화한 '인공지능 행동규범' 수립('24.10.)
(KT) 안전하고 신뢰할 수 있는 인공지능 서비스 개발을 위한 책임있는 인공지능 체계('Responsible AI프레임워크') 수립('24.10.)
(LG AI연구원) 인공지능 위험을 사전에 파악하고 해결하는 '인공지능 윤리영향평가' 전사업 적용('24.3~)

앞으로도 6개 기업은 자사가 제공하는 인공지능 제품·서비스와 관련된 책임을 깊이 인식하고, 향후 인공지능 신뢰·안전성 확보를 위한 노력을 더욱 강화해 나갈 것을 약속하였다.

③ '생성형 인공지능 잠재적 위험성.취약점 발굴 경진대회(AI레드팀 챌린지)'('24.5.) 결과 소개

올해 4월, 인공지능 연구자·대학생·일반시민 등 1,000여명이 참여하여 국내 인공지능 기업 4개사(네이버·SKT·업스테이지·포티투마루) 생성형 인공지능 모형(LLM)을 대상으로 잠재적 위험·취약점(유해정보, 편견·차별 등)을 발굴하는 '생성형 인공지능 취약 위험성 발굴 경진대회(레드팀 챌린지)' 행사를 개최한 바 있다.

이날 학술회의에서 한국정보통신기술협회(TTA)는 대규모 참여자의 공격 시도를 분석하여 식별된 7가지 분야* 주요위험과 다양한 공격기법(거부 무력화, 혼동유도)을 소개하였다.

* 잘못된 정보, 편견 및 차별, 불법 작품(콘텐츠) 생성 및 정보 제공, 탈옥, 사이버 공격, 개인의 권리 침해, 일관성

과기정통부는 취약점·위험성 발굴 경진대회(레드팀) 챌린지 결과를 활용하여 생성형 인공지능 안전 체계를 구축해 나가고, 국내기업이 자사 생성형 인공지능 모형의 잠재위험을 사전에 파악하여 신뢰·안전성을 확보할 수 있도록 지원할 계획이다.

④ '제2회 인공지능 신뢰성 대상' 시상

인공지능 신뢰성의 중요성과 인식을 업계 전반에 확산하고, 우수 인공지능 제품·서비스에 대한 홍보를 지원하기 위한 목적으로 '제2회 인공지능 신뢰성 대상' 시상식을 개최하였다.

인공지능 분야 산·학·연 전문가로 구성된 심사위원회(위원장 : 한국과학기술원 <KAIST> 오혜연 인공지능 연구원장)는 접수된 총 41개의 인공지능 제품·서비스를 대상으로 ①신뢰성·품질 이해 및 방침 적용수준, ②신뢰성·품질수준, ③신뢰성·품질 관리 우수성, ④시장성 등을 종합적으로 평가(1차:11.4., 2차:11.19.)한 결과, 총 8개의 인공지능 제품·서비스*가 선정되었다.

* △대상 1개(과기정통부 장관상, 상금 1,000만원), △최우수상 1개(과기정통부 장관상, 상금 700만원), △우수상 6개(한국정보통신기술협회 등 기관장상, 각 상금 300만원)

대상은 다비오가 개발한 ‘다비오 어스아이2.0*’가 수상했으며, ‘신뢰할 수 있는 인공지능 개발안내서’에 기반하여 자체적인 데이터 품질관리 절차를 수립·적용하여 높은 평가를 받았다.

* 위성, 영상 이미지의 공간정보 분석 등을 통해 시각화, 점검, 탐색 등 기능을 제공

< '제2회 인공지능 신뢰성 대상' 수상제품 >

상훈	회사명	제품명
대 상(장관상)	다비오	Dabeeo Eartheye 2.0
최우수상(장관상)	액션파워	daglo (다글로)
우수상(한국정보통신기술협회장상)	단감소프트	CareerRecommender v1.0
우수상(한국정보통신기술협회장상)	마크애니	SmartEYE for Finder v1.0
우수상(정보통신정책연구원장상)	셀렉트스타	LLM 신뢰성 평가 서비스
우수상(정보통신정책연구원장상)	튜닝	AI가드레일 솔루션패키지 v1.0
우수상(전자신문 사장상)	메디픽셀	MPXA-2000B v2.0.0
우수상(전자신문 사장상)	테스트웍스	블랙올리브 v1.0.0

송상훈 과기정통부 정보통신정책실장은 “정부가 지속적으로 추진해 온 인공지능 신뢰·안전성 정책의 노력에 호응하여, 최근 산업계·학계에서도 인공지능 신뢰·안전성 전담조직 설치와 투자 확대 등 자발적인 노력이 확대되고 있다.”고 강조하며,

“민간자율에 기반한 책임있는 인공지능 개발·활용 문화가 확산될 수 있도록 정책 지원을 강화하는 한편, 첨단 인공지능으로 인한 잠재적 위험에 대비하여 인공지능 안전연구소를 출범(11.27., 예정)하고 국가차원에서 체계적으로 대응해 나가겠다”고 강조하였다.

담당 부서	과학기술정보통신부 인공지능기반정책과	책임자	과 장	남철기 (044-202-6270)
		담당자	사무관	이재호 (044-202-6262)



더 아픈 환자께 양보해 주셔서 감사합니다

가벼운 증상은 동네 병·의원으로



제2회

인공지능

신뢰성

대상

2024. 11. 26. 화 10:00~17:00

피스앤피크컨벤션 3F 로얄홀 (서울 전쟁기념관)

AI

신뢰·안전성

컨퍼런스

10:00 ~ 10:15 (15분)	개회식	개회사 손승현 회장 한국정보통신기술협회, 김정연 원장직무대행 정보통신정책연구원 축사 송상운 정보통신정책실장 과학기술정보통신부
10:15 ~ 10:45 (30분)	시상식	제2회 인공지능 신뢰성 대상 시상식 및 기념촬영
10:45 ~ 11:00 (15분)	해외 기조 강연 (녹화 버전)	AI 안전의 방향 Yoshua Bengio 퀵백 AI연구소 교수, 국가AI위원회 글로벌 자문위원
11:00 ~ 11:30 (30분)	국내 기조 강연	글로벌 시생태계에서의 AI안전·신뢰 중요성 Alice Oh KAIST AI센터장 교수
11:30 ~ 13:00 (90분)	오찬	
제1세션 "위험"		
13:00 ~ 13:30 (30분)	발표	레드팀 챌린지 결과와 향후 계획 곽준호 팀장 한국정보통신기술협회 AI 신뢰·안전 확보를 위한 데이터셋 최호진 교수 KAIST
제2세션 "행동"		
13:30 ~ 15:30 (120분)	발표	AI 서울 정상회의, 그 후 1. 카카오 2. LGAI연구원 3. 네이버 4. KT 5. SKT * 상기 5개 기업은 "Seoul AI Business Pledge"에 서약하였으며, 이 중 네이버는 "Frontier Safety Commitment"에도 서약
15:30 ~ 15:40 (10분)	커피 브레이크	
제3세션 "기술"		
15:40 ~ 16:50 (70분)	발표	1. 안전 확보를 돕는 AI 모델 박규병 대표 휴넷 2. LLM 안전 및 신뢰 확보의 모범 사례 김동환 대표 포리투마루 3. LLM 신뢰성 평가 프로세스 및 Use Case 김세걸 대표 셀렉트스라
	토론	좌장 신준호 센터장 한국정보통신기술협회 패널 제3세션 발표자
16:50 ~ 17:00	폐회 및 장내 정리	