

보도시점 2025. 7. 8.(화) 15:00
(2025. 7. 9.(수) 조간)

배포 2025. 7. 8.(화) 07:30

인공지능(AI) 대전환 시대, 민·관이 함께 ‘신뢰’ 기반의 가드레일 구축 방안 모색

- 개인정보위, 「인공지능 프라이버시 민·관 정책협의회」 제4차 전체회의 개최
- ▲「인공지능 프라이버시 리스크 진단 및 인증 방안 연구」, ▲「생성형 인공지능(AI) 개발·활용을 위한 개인정보 처리 안내서(안)」의 주요 내용 공유 및 의견수렴

인공지능 대전환이라는 시대적 흐름에 발맞춰, 신기술 발전에 걸맞게 개인정보 보호 체계를 선진화하기 위해 민·관이 함께 머리를 맞댔다.

개인정보보호위원회(위원장 고학수, 이하 ‘개인정보위’)는 7월 8일(화) 오후 3시, 서울 서초구 소재 ‘서울 AI 허브’에서 민간 전문가 및 유관 부처 관계자 등이 참석한 가운데 「인공지능 프라이버시 민·관 정책협의회*」(이하 ‘정책협의회’) 제4차 전체회의(이하 ‘전체회의’)를 개최하고, 「인공지능 프라이버시 리스크 진단 및 인증 방안 연구」, 「생성형 인공지능(AI) 개발·활용을 위한 개인정보 처리 안내서(안)」에 대해 논의했다.

* 학계, 법조계, 산업계, 시민단체 등 인공지능 분야 차세대 전문가 30여 명으로 구성(23.10월~)

인공지능의 발전은 국민 일상에 편의성과 혁신을 안겨주고 있지만, 지속 가능한 발전이 되기 위해 기술로부터 발생할 수 있는 개인정보 침해 위험 등 잠재적 부작용을 예방하고 완화할 수 있는 제도적 안전장치가 필요하다.

개인정보위는 작년 12월 「인공지능(AI) 프라이버시 리스크 관리 모델」을 통해 인공지능 모델에 내재한 프라이버시 리스크 관리 절차와 리스크 유형 및 경감 조치 등을 안내한 바 있다. 다만, 이를 체계적으로 진단·측정할 수 있는 기준 제시는 후속 과제로 유보하여, 현재 국내 인공지능 환경에 적합한 사전적·예방적 관점의 구체적 안전관리 체계 구축 방안을 다방면에서 검토하고 있다.

이번 전체회의에서 김병필 카이스트 교수(정책협의회 1분과장)는 인공지능 프라이버시 리스크에 대한 구체적 진단 및 인증 방안에 대해 제언했다. 이 자리에서는 산업계·학계 중심으로 논의되는 프라이버시 리스크 진단의 최신 기술 및 우수 사례를 비교·분석한 결과와 인공지능에 내재된 프라이버시 리스크를 체계적으로 진단·평가할 수 있는 프레임워크 구축 방안 등에 대한 관점을 공유했다.

개인정보위는 이날 민간 전문가 및 관계부처 등에서 제시한 의견을 검토·반영하고, 후속 정책연구를 추진하여 향후 인공지능 프라이버시 리스크에 대한 체계적 진단·평가 체계 기반 마련을 위한 구체적인 방법론을 도출할 계획이다.

아울러, 「생성형 인공지능(AI) 개발·활용을 위한 개인정보 처리 안내서(안)」에 대해서도 논의했다. 최근 생성형 인공지능 기술이 급속히 발전하면서 일상과 사회 전반으로 빠르게 확산되고 있다. 특히, 공공민간의 많은 기관에서 생성형 인공지능을 도입해 각 분야에 적합한 인공지능 서비스를 개발·운영하는 등 인공지능 전환이 가속화되면서 개인정보 처리와 보호 관점에서 복잡한 법적·기술적 이슈*가 제기되고, 프라이버시 측면에서도 위험이 커지고 있다.

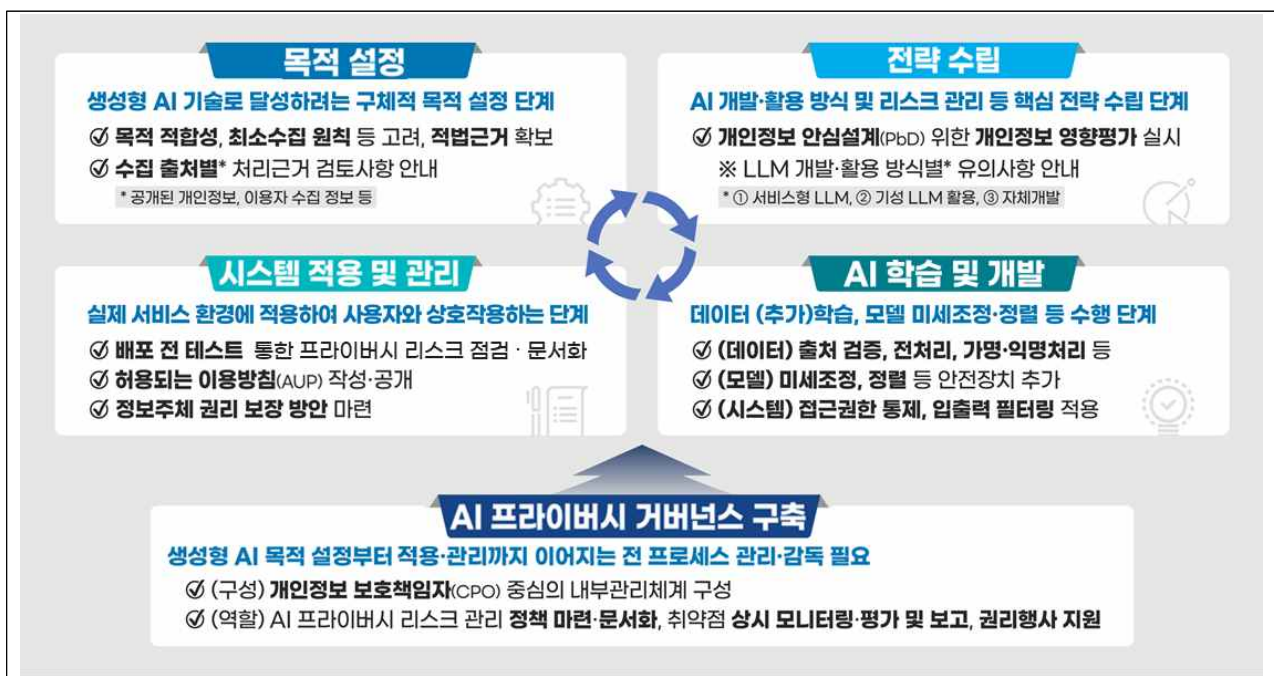
* 과거에는 정보주체의 동의나 계약 등에 따라 개인정보를 수집·이용했다면 인공지능 환경에서는 데이터 확보 경로와 처리 흐름이 복잡해져 개인정보 처리의 적법 근거가 불명확해지는 추세

이에 개인정보위는 인공지능 기술 및 서비스를 개발하거나 고도화하는 기관의 법적 불확실성을 해소하고, 개인정보 처리의 적법성, 안전성 등 개인정보 보호의 핵심 원칙을 확보하기 위해 생성형 인공지능 기술에 특화된 개인정보 보호 기준, 안전장치, 프라이버시 내부관리체계 등에 대한 명확한 기준을 마련해 왔다.

이번 안내서에는 생성형 인공지능의 수명주기 각 단계에서 개인정보 처리 및 보호를 위해 고려해야 할 적법성·안전성 확보 기준 등을 소개하여 현장의 법적 불확실성을 최소화하고, 생성형 인공지능 개발·활용에 개인정보 보호 관점이 체계적으로 반영될 수 있도록 했다.

특히, 이러한 내용들이 단순한 추상적 권고에 그치지 않고 실무 현장에 곧바로 적용될 수 있도록 ▲생성형 인공지능 관련 개인정보 보호법 주요 내용, ▲인공지능 개인정보 처리 안내서, ▲조사·처분 등 집행 사례, ▲규제샌드박스 등 혁신지원 사례를 포함한 그간의 정책적 경험을 활용하여, 인공지능 개발·구현 단계에 맞춰 도출한 개인정보 처리 및 보호 관련 고려 사항을 알기 쉽게 안내했다. 또한, 미국·영국·유럽연합(EU) 등 주요국 최신 정책자료를 참고하여 국제적 흐름도 반영했다.

<생성형 인공지능 개발·활용 단계별 고려 사항>



개인정보위는 이번 전체회의 논의를 통해 제시된 시민단체, 관계부처, 전문가들의 의견을 반영하여 「생성형 인공지능(AI) 개발·활용을 위한 개인정보 처리 안내서」를 7월 말까지 마련할 계획이다.

김병필 카이스트 교수는 “지금 인공지능은 멀티 에이전트 시스템으로 빠르게 나아가고 있다.”라며, “변화하는 환경 속에서 발전된 기술을 활용하면서도 충실하게 개인정보를 보호하기 위해 민·관이 함께 지혜를 모으는 것이 꼭 필요하다.”라고 밝혔다.

고학수 개인정보위 위원장은 “안전한 가드레일 구축은 신뢰할 수 있는 인공지능 발전의 선결 조건이며, 이러한 전제를 바탕으로 ‘인공지능 혁신을 위한 프라이버시’와 ‘프라이버시를 위한 인공지능 혁신’이라는 두 과제가 상호 발전적인 관계 속에서 기술 수용성과 사회적 신뢰를 높이는 기제가 될 수 있을 것”이라며,

“앞으로도 개인정보위는 인공지능 3대 강국 도약을 지원하기 위해 공공과 민간에 축적된 고가치 데이터가 안전하고 책임감 있게, 나아가 사회·경제적으로 ‘최적’ 활용될 수 있도록 합리적이고 유연한 정책 설계 노력을 지속하고 관련 기반을 마련해 나가겠다.”라고 밝혔다.

담당 부서	개인정보정책국 인공지능프라이버시팀	책임자	팀 장	구민주 (02-2100-3073)
		담당자	사무관 주무관	전이루 (02-2100-3078) 최홍석 (02-2100-3167)



□ 회의 개요

- (일시) '25. 7. 8.(화) 15:00 - 16:30
- (장소) 서울 AI 허브 메인센터, 3층 컨퍼런스룸A
※ 서울시 서초구 태봉로 108(우면동)
- (주요 참석자)
 - (정부) 개인정보위 위원장, 개인정보정책국장, AI프라이버시팀장, 과기정통부 인공지능기반정책관, 국가인공지능위원회 정책기획국장, 방통위 방송통신이용자정책국장, 행안부 공공지능데이터국장 등
※ 유관기관 의견수렴 등을 위해 AI안전연구소, TTA(한국정보통신기술협회) 등 관계자 참석
 - (민간) 1~3분과 분과장 및 민간위원 등
- (주요 내용) 논의 안전에 대한 보고 및 의견수렴

< 논의 안전 >

- ① 인공지능 프라이버시 리스크 진단 및 인증 방안 연구(김병필 카이스트 교수)
- ② 「생성형 인공지능(AI) 개발·활용을 위한 개인정보 처리 안내서」(인공지능프라이버시팀)

□ 행사 세부계획

시 간	주요 내용	비 고
15:00~15:01	개회(1')	* 사회: 총괄팀장
15:01~15:05	모두발언(4')	개인정보보호 위원장
15:05~15:17	(안전1) 「인공지능 프라이버시 리스크 진단 및 인증 방안 연구」 주요 내용(12')	김병필 교수 (KAIST)
15:17~15:29	(안전2) 「생성형 인공지능(AI) 개발·활용을 위한 개인정보 처리 안내서」 주요 내용(12')	담당 사무관
15:29~16:28	전체 토론(59')	* 사회: 총괄팀장
16:28~16:30	마무리 말씀(2')	개인정보보호 위원장

※ 현장 참석 기자단에게는 안전1 발표까지 공개