

2022년 제10호 (통권29호)



데이터산업 동향 이슈 브리프

ISSUE BRIEF

2022.10

데이터보호와 관련 기술 도입 논의 : 빅데이터·머신러닝을 중심으로

데이터보호와 관련 기술 도입 논의 : 빅데이터·머신러닝을 중심으로

I. 서론	1
II. 빅데이터와 데이터보호	4
III. 개인정보 이용 고지 및 동의 모델의 문제	8
IV. 인공지능 기반 서비스 운영에서의 개인정보보호	10
V. 사후 데이터 관리 책임	14
VI. 리스크 관리와 데이터 윤리 기반 프레임워크 구축	17
VII. 결론	19

요 약

- 일상의 금융서비스 부문에서 빅데이터와 인공지능 기술 응용이 점차 증가하면서 편의성과 금융 포용성이 높아지는 장점이 있으나, 해당 기술을 통한 추론의 편향성과 기술 응용 방식에 대한 일반적 설명이 어려운 점 등의 문제가 드러나고 있으며, 소비자데이터 보호 및 개인정보보호 측면에서의 정확성, 공정성, 윤리성 등의 분야에서 문제될 수 있는 부분들을 새로운 프레임워크에서 논의 할 필요가 있어 본 보고서에서는 여러 국가 및 관계 기관들이 취하고 있는 접근과 향후 고려해야 할 사항들에 대해 검토함
- 본 보고서에는 빅데이터와 머신러닝 전반에 적용할 수 있는 표준 및 기술을 개발하고, 데이터 프라이버시 보호와 혁신 간의 균형을 맞추기 위해 하기와 같은 내용을 다루고 있음
 - 개인정보의 이용 및 공유에 대해 소비자에게 알기 쉽게 설명을 제공하는 등 소비자 공지 방식 개선, 소비자 개인정보보호를 강화하는 데 필요한 조치 마련
 - 무분별한 소비자 데이터 공유를 막기 위해 개인정보 사용 및 공유 관련 규제 강화
 - 인공지능 및 머신러닝 모델 설계에 프라이버시 원칙을 적용하기 위한 표준 개발
 - 인공지능 컴퓨터 프로그래밍에 대한 윤리적 표준 개발, 신규 문제 식별 및 접근 방식 개발
 - 데이터를 활용한 추론 과정에서의 신뢰도 확립을 위한 표준 개발
 - 출력 데이터의 평가와 머신러닝 모델의 개인정보보호 및 차별 방지 원칙에 대한 준수 여부 확인
 - 머신러닝 시스템이 도출한 추론에 사용된 데이터 정보 및 추론 프로세스를 설명하기 위한 표준 개발
 - 자동 결정에 대한 소비자 이의 제기 절차 마련 및 인공지능 및 머신러닝 모델의 설계 및 운영에 대한 책임 평가를 위한 프레임워크 마련
- 우리나라 정부도 데이터의 안전한 활용을 위한 및 신기술 개발 및 적용과 법적 제도 정비를 추진하고 있음
 - 개인정보보호위원회는 2021년 11월 데이터 경제 시대에 정보주체인 소비자의 권리를 보호하고 안전한 활용을 지원하기 위해 개인정보에 특화된 보호·활용 기술에 대한 연구개발을 본격 추진한다고 발표함
 - (비정형 데이터에서 개인정보 탐지 기술) 대화형 텍스트, 음성 및 영상 데이터 등 비정형 데이터에 대한 개인정보 식별을 위해 필요한 기술 개발
 - (가명·익명정보 안전성 평가 기술) 데이터 유형별 가명·익명처리 기술의 재식별 가능성 및 공격모델 등을 연구하여, 기술의 안전성 검증
 - (프라이버시 보존형 개인 맞춤 서비스 기술) 개인정보를 노출하지 않고도 특정 요건·자격 유무를 증명할 수 있는 고속 영지식 증명 생성기술 개발
 - (개인정보 동의 관리 기술) 개인정보 수집, 제3자 제공에 대한 동의, 동의 철회 등 정보 주체의 자기정보 통제권 보장을 위한 UX 중심의 설계 방안 마련
 - 올해 출범한 디지털플랫폼정부도 7월 새로 데이터보안 기술을 도입할 것이라 밝힘
 - 한국인터넷진흥원은 디지털플랫폼 정부에 망분리 및 클라우드 보안인증 제도를 혁신적으로 개선하고, 제로 트러스트, 블록체인, 인공지능 등 최신 보안 기술을 도입할 것이라고 밝힘

PART I

서론

▶ 빅데이터, 인공지능 및 머신러닝이 일상의 금융 서비스, 알고리즘 등에 많이 활용되면서 이에 대한 관심이 커지고 있음

- 특히 금융 서비스 제공 기업에서 개인 식별 가능 데이터¹⁾를 이용해 신용도 평가 및 대출 실행 관련 결정을 수행하고 있음
 - 많은 금융 서비스는 위험 평가와 관리에 의존하기 때문에 대출자에 대한 데이터를 최대한 많이 수집하여 신용도를 평가하고자 하며, 빅데이터를 통해 과거에 채무 상환과 채무 불이행 이력을 확인하고 신용도를 도출할 수 있음
 - 빅데이터 분석은 특히 과거 은행거래 내역을 중심으로 판단했던 신용평가 방식을 보다 다양한 금융거래 및 소비 정보 등 더 많은 정보를 바탕으로 하는 신용평가 방식으로 개선할 수 있으며, 미국의 3대 신용 정보기관인 에퀴팩스(Equifax)는 빅데이터 분석을 사용하여 신용 평가 모델의 예측 능력을 크게 향상시켰다고 밝힘
 - 디지털 금융 서비스 제공자는 빅데이터를 활용하여 이처럼 상업적 이익을 창출할 수 있을 뿐만 아니라 소비자의 배경과 관심사에 대한 정보 분석을 통해 금융 서비스에 대한 접근성을 개선하여 사회에 긍정적 외부효과를 제공할 수 있음
- 빅데이터를 통한 신용 평가는 기존 신용 평가 모델을 혁신적으로 변화시킴
 - 예를 들어, 핀테크 기업인 업스타트(Upstart)는 대출 과정을 자동화하였으며 이 과정에서 대출자의 교육 이력, 시험 점수, 학업 분야 및 직업 이력 데이터를 기반으로 신용도를 평가함
- 인공지능과 머신러닝은 신용평가뿐만 아니라 보험 등 리스크 평가와 관련된 분야에서 많이 활용되고 있음에서도 활용되고 있음

▶ 이처럼 인공지능과 머신러닝에 의존하는 서비스 이용이 증가하면서 데이터보호에 대한 중요성 또한 높아져 각국의 규제 당국에서 데이터 보호를 위한 기술 및 표준을 강화해야한다는 목소리가 높아지고 있음

- 세계은행(World Bank)은 2018 G20 정상회의 이전에 ‘공식 경제의 신용평가 제고를 통해 개인과 중소기업의 디지털 금융서비스 접근성을 확대하기 위한 대체 데이터 활용 보고서(Use of Alternative Data to Enhance Credit Reporting to Enable Access to Digital Financial Services by Individuals and SMEs operating in the Informal Economy)’를 발표함²⁾
 - 해당 보고서는 대체 데이터 활용에 대한 데이터보호 문제를 분석하고 정책 입안자와 규제 당국에

1) 개인 식별 가능 데이터는 개인의 인터넷 거래 및 사용 내역, 공공 및 민간단체 가입 여부, 소셜미디어 이용 내역 등 정보주체의 식별이 가능한 정보를 포함하고 있으며, 기업과 정부는 이러한 데이터를 수집, 처리 및 제3자와 공유함

2) GPFI, Use of Alternative Data to Enhance Credit Reporting to Enable Access to Digital Financial Services by Individuals and SMEs operating in the Informal Economy, Guidance Note, 2018

권고 사항을 제시함

- 개인정보보호법의 세 가지 핵심 원칙은 목적 명시(Purpose Specification), 데이터 최소화(Data Minimization), 특정 집단(인종, 성별, 종교 및 기타 집단)의 데이터 처리(The Treatment of data of special categories of groups)인데, 인공지능 및 머신러닝이 이러한 원칙을 위협하는 부분이 있음

▶ 이러한 상황에서 우리나라 정부도 데이터의 안전한 활용을 위한 및 신기술 적용과 법적 제도 정비를 추진하고 있음

- 개인정보보호위원회는 데이터 경제 시대에 정보주체인 소비자의 권리를 보호하고 안전한 활용을 지원하기 위해 개인정보에 특화된 보호·활용 기술에 대한 연구개발을 본격 추진한다고 발표함³⁾
 - 이를 위한 첫 단계로 향후 5년간의 연구개발 방향을 담은 개인정보보호·활용 기술 R&D 로드맵('22~'26)을 공개함
 - 해당 로드맵은 개인정보보호위원회가 마련한 개인정보보호·활용 기술분류체계와 국내·외 기술 및 표준화 동향 등을 고려해 수집됐으며, 정보주체 권리보장, 유·노출 최소화, 안전한 활용 등 3대 분야 11대 핵심기술과 37개 세부기술을 담음
- 한국인터넷진흥원의 경우 디지털플랫폼정부 보안 방향과 전망에 대해 의견을 제시함⁴⁾
 - 한국인터넷진흥원은 현재 추진 중인 디지털플랫폼정부가 모든 데이터가 연결되는 플랫폼이기 때문에 사이버 위협을 방지하기 위한 다양한 혁신 개선 및 보안 기법 도입을 계획함
 - 따라서, 한국인터넷진흥원은 망분리 및 클라우드 보안인증 제도를 혁신적으로 개선하고, 제로 트러스트, 블록체인, 인공지능 등 최신 보안 기법을 도입하고 확산시키겠다고 공언함
 - 아울러, 개인정보 활용 과정에서 이상행위 탐지 및 재식별 방지 체계를 마련하고 마이데이터 전 과정의 보안과 감독을 강화하겠다고 강조함

▶ 본 보고서(*)에서는 데이터보호제도 및 기술에 대한 필요성과 이슈 사항들에 대해 다룸

(*) Big data, machine learning, consumer protection and privacy

- 데이터 분야의 혁신 및 경제 생산성 제고가 가져다 줄 혜택과 소비자데이터 보호와 개인정보보호 정책의 목적 달성 간의 균형을 어떻게 조절해 나갈 것인가 하는 바를 두고 정책 입안자, 규제기관, 디지털 금융 서비스 제공업체 그리고 시장 참여자간의 컨센서스는 부재한 상황이며 일의적으로 적용가능한 규제가 없어 관계자 간의 담론이 필요함
- 이에 본 보고서에서는 관련 기업, 학계, 정부 등이 소비자데이터 보호 및 개인정보보호를 위해 취하고 있는 접근과 제안 사항을 검토함
- 특히 2장과 3장에서는 빅데이터와 머신 러닝, 관련 과제 등 기초 사항을 다루며, 4장에서는 빅데이터와 머신 러닝의 기술과 시장 동향, 사용되는 데이터의 종류, 이러한 데이터를 사용하는 프로파일링 및 자동화된 결정을 시작으로 주요 개념을 소개함

3) 권정수, “개인정보보호·활용 기술 R&D 로드맵 공개…11대 핵심기술 개발 추진”, IT Daily, 2021년 11월 10일

4) 김가은, “‘디지털 플랫폼 정부’는 AI가 지킨다…최신 보안기법 도입으로 ‘新보안체계’ 구축”, TechM, 2022년 7월 14일

- 5장에서는 주로 개인 데이터를 수집, 사용 및 제3자에게 이전하는 방법과 목적에 대해 소비자에게 어떻게 통지해야 하는지, 그리고 개인 데이터의 사용을 정당화하기 위한 소비자 동의를 얻기 위한 요건에 대해 논의함
- 6장에서는 머신러닝 모델의 정확성, 편향 및 차별적 처리, 데이터 침해 및 재식별을 포함하여 기업이 개인 데이터에 대해 수행할 수 있는 책임과 관련한 사항을 논의함
- 7장에서는 빅데이터 및 머신러닝을 활용하는 기업의 개인정보보호법 위반에 대한 책임을 물을 수 있는 소비자의 권리를 다룸
 - 소비자는 자신의 데이터에 대한 접근권, 오류를 수정하고 삭제 요청을 할 수 있는 권리, 복잡한 기계 학습 모델 출력에 대한 투명성 요구, 의사결정에 이의를 제기하고 수정을 요구할 권리가 있음

PART II

빅데이터와 데이터보호

▶ 빅데이터와 인공지능의 관계는 양방향으로 볼 수 있음

- 빅데이터(Big Data)는 디지털 환경에서 생성되는 데이터로, 그 규모가 방대하고 생성 주기도 짧으며 형태도 수치 데이터 뿐만 아니라 문자와 영상 데이터를 포함하는 대규모 데이터를 칭함
 - 빅데이터는 소비자의 통신, 거래 및 움직임을 기록하는 수많은 애플리케이션과 센서를 이용하여 수집되며, 분산 데이터베이스는 데이터를 저장하고 고속 통신은 데이터를 고속으로 전송하여 데이터 분석 비용을 절감함
- 인공지능(Artificial Intelligence)은 학습능력과 추론능력, 지각능력, 자연언어의 이해능력 등을 컴퓨터 프로그램으로 실현한 기술로, 인공지능의 한 형태인 머신러닝(Machine Learning)은 대규모 데이터 세트의 패턴을 인식하여 여러 분석 계층에서 성능을 향상시키는 시스템의 능력을 의미함
 - 머신러닝 알고리즘은 미리 프로그래밍된 로직만 따르는 것이 아니라 과거 데이터를 기반으로 모델을 구축하여 예측 또는 결정을 도출함
- 빅데이터는 인공지능과 머신러닝에 의존해 빅데이터 데이터 세트에서 가치를 추출하고, 머신러닝은 방대한 양의 빅데이터에 의존해 학습함

▶ 과거 금융 서비스에서 사용하는 자료에는 지역 은행 관리자 및 보험중개인, 신용 조회국 자료 정도가 포함되었으나, 오늘날에는 광범위한 디지털 출처를 통해 수집된 빅데이터를 포함함

- 금융 서비스에서 사용되는 대체 데이터는 대부분 통신 네트워크 운영자의 서비스에서 파생됨
 - 통신사는 일반적으로 고객의 전화 통화 내용에 대한 데이터를 수집하고 사용할 수 없도록 제약을 받지만 통화기록(Call Detail Record) 등 메타데이터에 접근할 수 있음
 - 메타데이터는 고객의 통신 서비스 사용에 대한 데이터로서, 누가 언제, 얼마나 오랫동안 누구와 통신했는지와 통화가 이루어진 위치 등을 포함함
 - 통신사는 또한 모바일 결제 서비스를 운영하고 있기 때문에 고객의 결제액, 잔고, 송수신처, 부채 상환 내역 등에 대한 데이터를 보유함
 - 이러한 데이터를 조합하면 개인 간의 관계와 현금 흐름을 유추할 수 있기 때문에 통신사는 신용점수를 산정할 수 있도록 금융사와 알고리즘을 적용한 모바일 머니 사용 데이터 및 통화기록 데이터를 공유함
- 이 때문에 많은 국가의 전기통신법 및 자격법에서는 통신사가 전자통신 서비스를 사용하여 전송된 통신 내역을 활용하거나 공개 또는 기록하는 것을 명시적으로 금지하는 조항을 포함하고 있으며 메타데이터도 규제 사항으로 포함하는 추세임

- 예로 EU의 프라이버시 지침(ePrivacy Directive)은 전자 통신 서비스 규제를 구체화하는 프라이버시 규정(ePrivacy Regulation)으로 변경되었으며, 프라이버시 규정은 통화 세부 내역 및 메타데이터 등을 모두 다루고 있음
- 그러나, 아직까지는 메타데이터를 규제하는 법률이나 기술 표준이 보편적으로 도입되지 않았으며, 메타데이터 규제가 시행이 되더라도 고객의 동의만 있으면 통신사에서 메타데이터를 활용할 수 있음
- 통신사 서비스 외에 웹 브라우저 및 휴대폰 애플리케이션에서 파생된 대량의 데이터 역시 수집 및 공유됨
 - 데이터 시장은 웹 추적 뿐만 아니라 사용자가 스마트폰에서 컴퓨터 및 태블릿과 연동해 서비스를 이용하는 것까지 데이터를 수집하기 때문에 사물인터넷(IoT)이 발전함에 따라 데이터 수집량도 늘어날 것으로 예상됨
 - 이렇게 광범위한 데이터 출처에서 발생된 브라우저 및 검색 기록, 소셜 네트워크, 스트리밍한 비디오 및 음악, 소매 구매 기록, 블로그 게시물 및 온라인 리뷰 등 다양한 데이터를 통해 사용자의 위치나 선호 등을 추적할 수 있음
- 이와 더불어 빅데이터를 생성하기 위해 사용되는 대체 데이터⁵⁾의 출처는 매우 다양함
 - 금융 통합업체의 데이터(Data from financial aggregators), 신용카드 데이터(Credit card data), 지리공간 및 위치 데이터(Geospatial and location data), 웹 스크래핑 데이터 세트(Web scraping datasets), 앱 참여 데이터(App engagement data), 미국 세관 발송 데이터(Shipping data from U.S. customs), 광고비 지출 데이터(Ad spend data), 응용 프로그램을 통한 데이터(Data made available through APIs), 센서 및 라우터를 통해 생성된 위치/유동 인구 데이터(Location/foot traffic data from sensors and routers), 소셜 미디어 데이터(Social media data), 공급망에서 취득한 B2B 데이터(B2B data acquired from parties in the supply chain), 농업 데이터(Agriculture data), POS 데이터(Point of sale data), 의약품 처방 자료(Pharmaceutical prescription data)가 있음
- 빅데이터와 머신러닝, 인공지능은 프로파일링과 자동 결정을 통해 사회적으로 이익을 발생시킴
 - 프로파일링(Profiling)은 개인의 관심사, 개인적 선호, 행동, 업무 수행, 경제 상황, 건강, 신뢰성, 위치 또는 움직임 등을 평가, 분석 또는 예측하기 위해 개인 데이터를 자동 처리하는 행위를 뜻하며, 데이터 분석을 통해 개인과 개인이 속한 집단 구성 간의 연결 고리를 식별할 수 있음
 - 프로파일링을 통한 추론과 예측은 표적 광고에 사용될 수 있고, 자동화된 결정을 내리는 데 사용될 수 있음
 - 자동화된 결정(Automated decision)은 머신러닝과 빅데이터를 활용한 프로파일링을 통해 내린 추론과 예측을 기반으로 인간의 개입 없이 컴퓨터 처리 시스템에 의해 자동으로 결정을 내리는 것을 의미하며, 추론과 예측을 함으로써 기업에서 소비자의 선호와 요구에 맞는 제품과 서비스를 제공할 수 있도록 함
- 금융 서비스 산업의 여러 부문에서 빅데이터 및 머신러닝 응용프로그램이 도입되고 있음

5)

- 대출 및 보험 위험 평가, 투자 포트폴리오 관리, 로보어드바이저, 헤지펀드 및 금융 기관의 초단타매매(HFT)⁶⁾, 자산 관리, 유동성 및 외화 위험 및 금융시스템 위기관리 테스트(Stress test), 기업의 부정 행위 탐지, 보안 및 디지털 식별, 뉴스 분석, 고객 판매 및 추천, 고객 서비스와 같은 다양한 서비스에 응용프로그램이 도입되고 있음
- 일부 국가에서는 인공지능의 사용을 명시적으로 허가하는 법률을 제정하여 이러한 응용프로그램의 활용을 법적으로 허용하고 있음
- 멕시코는 2018년 기준 증권 시장법(The Securities Market Law)을 개정하여 자동화된 자문 및 투자 관리 서비스를 허용한 바 있음

▶ 기업과 소비자 간에 정보 비대칭성이 존재하기 때문에 국가에서 법률과 규제를 통해 소비자를 보호할 의무가 있음

- 기본적으로 기업은 소비자보다 더 많은 데이터와 지식을 가지고 있기 때문에 문제 발생 시 협상력의 심각한 비대칭성이 발생할 수 있음
 - 특히 오늘날 디지털 서비스가 대부분 인공지능 및 머신러닝과 같은 컴퓨터 프로그램에 의존해 제공되기 때문에 충분한 이해가 부족한 소비자와 기업 간의 신뢰가 훼손되고 디지털 서비스의 성장 전망을 저해할 수 있는 리스크가 있음
 - 따라서 국가는 공정성, 책임성 및 투명성(FAT, Fairness·Accountability·Transparency)의 가치를 보호하기 위해 소비자 데이터보호법을 제정함
- 소비자 데이터보호법은 일반적으로 소비자가 구매하는 제품 및 서비스와 관련된 특정 권리를 다루며, 이러한 권리에 다음이 포함됨
 - 제공된 제품 또는 서비스에 대한 정보 제공 등 구매 전 권리(사전 제공), 제품 또는 서비스 자체의 품질 및 기능(계약), 제공자에게 책임을 묻는 사후 관리 수단(사후 관리)
 - 이러한 FAT 가치는 사전 제공 및 계약 뿐만 아니라 사후 관리 단계에서도 적용이 되며, 소비자가 제품 또는 서비스에 대한 이의를 제기할 수 있도록 보호를 제공함
 - 아울러 소비자 데이터보호법은 기업과 소비자 간의 협상 여부와 상관 없이 잘못된 제품 설명, 불공정한 계약 조건, 결함 있는 제품, 보상 체계 부재로부터 소비자를 보호함

▶ 모든 빅데이터와 머신러닝 기술이 개인 데이터에 의존하거나 데이터보호 문제를 일으키는 것은 아니며, 데이터의 본질을 구분지을 필요가 있음

- 데이터를 통해 개인의 식별이 가능할 경우 해당 데이터는 개인 데이터에 해당하며, 개인 데이터를 사용하는 경우 프라이버시 침해의 우려가 발생함
 - 프라이버시는 개인의 개성, 자율성 및 존엄성 가치를 포함하는 개념으로, 디지털 환경에서 프라이버시를 보호하기 위해서 개인 데이터의 수집, 사용 및 공유에 대한 통제가 적용됨
 - 개인 데이터는 개인이 제공하는 데이터(사용자 이름 또는 우편번호 등)와 개인의 관찰 데이터(위치 데이터 등), 도출된 데이터(우편번호와 거주 국가 등), 추론 가능한 데이터(신용 점수 등)가 해당됨

6) 고성능 컴퓨터의 월등한 속도를 무기로 극초단타 거래를 해 수익을 쟁기는 시스템

- 빅데이터의 경우 개인 데이터가 익명화(Anonymization)⁷⁾ 또는 데이터 가명화(Pseudonymization)되었더라도 개인 재식별(Re-identification)이 가능한 경우 개인정보보호 위험이 발생함
- 이에 따라 점점 더 많은 국가에서 데이터 수집, 사용 및 공유의 최소화를 통해 개인 데이터와 개인정보를 보호하기 위한 법률을 제정하고 있음
- 데이터보호 대책으로는 수집되는 개인 데이터를 알 수 있는 소비자 권리, 부정확한 개인 데이터를 수정하고 불완전한 개인 데이터를 완성시킬 수 있는 권리, 개인 데이터를 삭제 받을 수 있는 권리, 개인 데이터를 제3자에게 이전할 수 있는 권리, 개인 데이터 처리에 이의를 제기할 수 있는 권리 보장 등이 있음
- 이러한 개인 데이터에 대한 권리 보장을 통해 소비자 권리를 보호하고 프라이버시 침해를 막을 수 있음

7) 데이터를 가공하여 본래의 개인정보를 알아 볼수 없도록 만드는 것

PART III

개인정보 이용 고지 및 동의 모델의 문제

- ▶ 개인정보보호법은 기업이 개인 데이터의 수집, 사용 및 공유 목적을 소비자에게 통지하고 동의를 얻도록 강제함
 - 개인정보보호법의 두 가지 핵심은 목적 명시와 데이터 최소화에 해당함
 - 기업은 데이터를 수집, 사용 및 공유하는 목적을 명시해야 하며, 명시한 목적에 따라 필요한 수준만큼의 데이터 수집, 사용 및 공유가 허용됨
 - 이는 원래 한 목적으로 수집된 데이터가 다른 목적으로 사용되는 기능 확대(Function Creep)가 발생하지 않도록 법에서 제한하기 위해서임
 - 그럼에도 불구하고, 데이터의 초기 수집 목적을 벗어나도 사용을 허용하는 사례가 때때로 존재함
 - 데이터가 통계적 목적이거나 과학적 연구 목적으로 사용되는 사례와 금융 서비스 제공을 위한 제품 개발에서 사용되는 사례 등 법률에서 개입하기 애매모호한 부분이 있음
 - 소비자는 자신의 개인 데이터가 수집, 사용 및 공유되는 것을 원하지 않는 경우 동의를 철회하겠다고 기업에 통지할 수 있으며, 소비자가 그러한 선택권이 없는 경우 기업에서 개인 데이터 사용 및 공유 목적을 명시해야 함
- ▶ 빅데이터 및 머신러닝 기술의 도입에 따라 개인정보보호법에서 강제하는 동의 모델링에 문제가 발생함
 - 머신러닝 기술을 이용하는 기업이 이러한 통지 및 동의 모델의 요구 사항을 준수하려면 개인 데이터 수집 목적을 상세화하고 목적에서 벗어난 사용을 막기 위해 운영을 면밀히 모니터링해야 함
 - 머신러닝은 본질적으로 데이터 패턴을 감지한 다음 더 깊은 층(Layer)을 파고들어 추가 패턴을 식별하기 때문에 원래 목적과 관련이 없는 사용 사례가 나타날 수 있음
 - 또한 머신러닝 기술은 시간이 지남에 따라 더 큰 데이터 세트에서 패턴을 더 효과적으로 탐지해낼 수 있기 때문에 기업에서는 최대한 많은 양의 빅데이터를 수집하고 오랫동안 유지하려고 함
 - 따라서 머신러닝 기술을 이용하는 기업은 데이터 최소화 개념에 반하여 기술을 활용할 가능성이 있고, 일일이 소비자에게 이를 통지하고 규정을 준수하는 것이 어려움
 - 이와 같은 어려움 때문에 데이터 사용 목적을 폭넓게 설명하는 것 또한 법적으로 허용되지 않을 수 있음
 - 소비자 통지 및 동의를 더욱 원활히 하기 위해 동의 사항을 더 잘 관리할 수 있는 기술과 서비스 개발 노력이 병행되고 있음

- 디지털 권리 관리와 개인 데이터 권한 부여, 개인 데이터 수집, 사용 및 공유에 관한 양측 자동 협상 등 여러 접근법이 구상 중에 있음
- 아울러 소비자가 개인 데이터 공유 선호도를 미리 설정하여 데이터 대리인이 권한을 위임받아 동의를 하는 방식도 있음
- 데이터 대리인은 데이터 수집 기업이 제공하는 소비자 교육 및 통지를 대신 받고, 데이터가 승인되지 않은 방식으로 사용되고 있는 경우 소비자에게 알리는 책임을 수행함

PART IV

인공지능 기반 서비스 운영 상 소비자데이터 보호의 한계

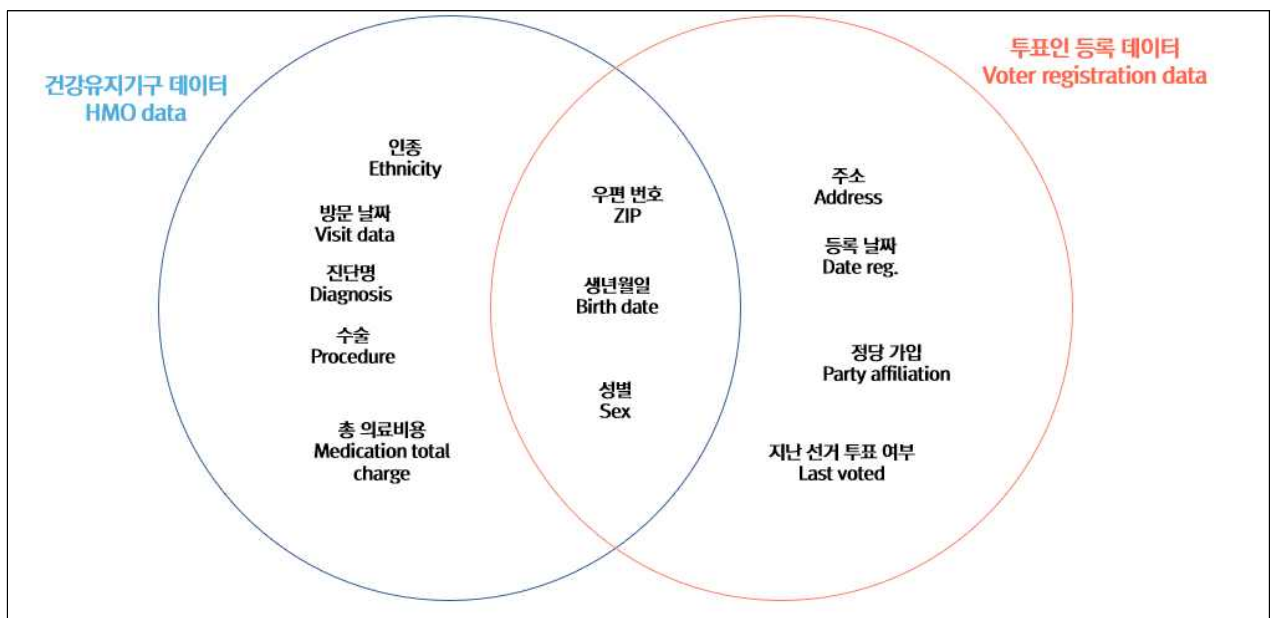
- ▶ 머신러닝 시스템은 입력 데이터의 정확도에 따라서 출력 데이터의 품질이 크게 차이날 수 있음
 - 머신러닝 시스템에 입력되는 방대한 양의 데이터 중에서는 구조화된 데이터가 있고 구조화되지 않은 데이터가 있기 때문에 정확도에 영향이 있을 수 있음
 - 데이터는 시간이 지남에 따라 출처가 다양해지고 출처와 수집 방식이 다양해지고 수집의 직접성이 떨어지는 경우가 있기 때문에 데이터의 신뢰도가 떨어지고 체계적인 업데이트 프로세스가 적용되지 않을 수 있음
 - 이러한 요인으로 인해 소비자에 대한 잘못된 정보가 생성될 수 있음
 - 따라서 개인정보보호법은 기업이 보유하고 처리하는 데이터의 정확성을 보장하기 위해 기업에 법적 책임을 부과함
 - 멕시코의 데이터보호 법률은 기업이 데이터베이스의 개인 데이터가 수집 목적에 맞게 올바르게 사용되고 있는지와 주기적으로 갱신되고 있는지 검증하도록 요구함
 - 금융 서비스 관련 법 또한 금융 서비스 내에서 사용되는 데이터의 정확성을 보장하도록 요구함
 - 그러나, 빅데이터와 머신러닝의 등장으로 기존 법률과 정책 지침이 현재 기술 발전을 따라잡지 못하고 있는 상황임
 - 빅데이터와 머신러닝은 당초 합의된 목적을 벗어난 개인 데이터를 수집해 사용할 수 있기 때문에 리스크를 내포하고 있음
 - 따라서, 많은 국가에서는 공정하고 정확한 신용평가(Credit Reporting)가 갖는 공공의 이익을 인식하여 신용 거래, 보험, 인허가, 소비자 주도 사업 거래, 고용 등을 위해 데이터를 활용하는 기업을 규제하고자 함
 - 이러한 기업들은 소비자의 데이터를 대량으로 다루기 때문에 데이터 정확성을 보장해야 하며, 공개 및 보고 의무에 의해 규제를 받음
- ▶ 기존 개인정보보호법은 현재 빅데이터 기반의 머신러닝 응용 결과 데이터까지 규제하고 있지 않는 한계가 있으며 그 결과 또한 부정확성과 편향성의 이슈가 발생할 수 있음
 - 개인 데이터의 수집, 사용 및 공유를 규제하는 데이터보호법은 빅데이터를 기반으로 한 머신러닝 모델의 결과 데이터까지는 규제하지 않음
 - 개인정보보호법은 차별을 방지하는 주목적이 있으며, 디지털 금융 포용 원칙(The High Level Principles for Digital Financial Inclusion)에서는 데이터가 디지털 금융 서비스와 관련하여 불공정하며 차별적인 방식으로 사용되어서는 안 된다고 명시하고 있음
 - EU의 일반데이터보호규정은 더 엄격한 제한을 위해 개인 데이터를 특별 범주로 나누어서 규제함

- 특별 범주는 인종적 기원, 정치적 의견, 종교적 또는 철학적 신념, 노동 조합 가입 여부, 유전자 데이터, 생체 인식 데이터, 성생활 또는 성적 지향에 관한 데이터에 해당함
- 머신러닝 알고리즘이 역사적 사례를 기반으로 하는 입력 데이터를 처리하는 경우 역사적으로 불리한 특정 모집단에 불리한 결과를 초래하는 문제가 있을 수 있음
 - 머신러닝 알고리즘은 과거의 차별적 정보를 그대로 반영하여 기존 편견을 악화시킬 수 있음
 - 한 집단의 은행 거래 내역이 다른 집단의 은행 거래 내역보다 적은 경우 머신러닝 알고리즘은 은행 거래 내역이 적은 집단에 불리한 결과를 내놓을 수 있음
- 편향적으로 결과를 도출할 수 있는 요소를 제거하는 기술을 적용한다고 해도 완벽하게 편향성을 해결할 수 없기 때문에 편향성의 영향을 평가하기 위한 시험이 진행 중임
- 편향성을 줄이기 위해 대표적으로 활용할 수 있는 방식에는 두 가지가 있음
 - 무작위성을 활용해 데이터의 편향성을 줄일 수 있는데, 예를 들어 A그룹에 대한 데이터로 결론을 도출하다 보면, 지역이나 연령의 편향성이 나타날 수 있기 때문에 다른 지역이나 연령이 분포된 B그룹 또는 C그룹의 데이터를 무작위로 투입해 A그룹 데이터가 가진 편향적 추정을 줄일 수 있음
 - 머신러닝 알고리즘에 투입되는 입력 데이터를 선택적으로 수정하여 출력되는 데이터의 편향성이 낮아지도록 개입하는 방안이 있음
- 이렇게 편향성을 줄이면 시간이 지남에 따라 알고리즘이 개선될 수 있으나, 현재 새로운 머신러닝 모델과 서비스가 지속적으로 시장에 나타나고 있기 때문에 편향성을 확인할 책임과 비용을 어느 주체가 부담해야 하는지에 대한 문제가 제시됨
 - 결론적으로, 자동화된 결정에 따라 편향적인 결과가 지속적으로 나오게 된다면 차별금지법을 위반하게 되므로, 기업에서 데이터보호 영향 평가(DPIA, Data Protection Impact Assessment)를 수행하고 윤리적 프레임워크를 도입해 편향성을 줄이도록 국가에서 모니터링해야 할 책임이 있음
- ▶ 아울러, 빅데이터와 머신러닝 모델은 ‘차등 가격’이라는 요소를 통해 불평등을 심화시킬 수 있음
 - 금융 서비스 제공자는 소비자보다 더 많은 정보를 보유하고 있기 때문에 이 정보 비대칭성을 이용해 동일한 제품에 대해 소비자에 따라 다른 가격을 부과할 수 있으며, 이러한 현상을 ‘차등 가격(Differential pricing)’이라고 일컫음
 - 예를 들어, 보험 회사는 교통 사고를 일으킨 적이 없는 개인에게 보험료를 인하해줌으로써 위험한 행동을 억제하고 경제적 효율성을 향상시킬 수 있음
 - 그러나 개인의 통제를 벗어나 발생할 수 있는 질병의 경우에 건강보험 가입자에게 불이익을 주기 때문에 불공정성을 초래할 수 있음
 - 빅데이터와 머신러닝은 개인 데이터에서 서비스에 대한 개인의 요구, 지불 능력 및 가격 민감도에 대한 추론을 이끌어냄으로써 차등 가격 책정에 관여함
 - 통상 서비스 제공자가 빅데이터를 보유하고 있기 때문에 서비스 제공자가 저소득자, 채무자 등 취약 계층 등에 차별적으로 가격을 제시할 수 있음
 - 이 때문에 빅데이터를 기반으로 하는 머신러닝 모델은 궁극적으로 계층별로 불평등을 심화시킬

수 있음

- ▶ 앞서 언급한 바와 같이 빅데이터 활용으로 인해 데이터 보안 침해 리스크가 높아지고 있으며, 편향성을 줄이고 보안 수준을 높이기 위해 여러 대책이 강구되고 있음
- 비식별화(De-identification)는 데이터 세트의 정보를 직접 식별하거나 간접적으로 식별하는 것을 막기 위해 장벽을 도입하는 것을 의미함
 - 정보 직접 식별은 추가 정보 없이 공용 도메인의 정보(사람의 이름, 전화 번호, 전자 메일 주소, 사진, 주민등록번호 또는 생체 인식 정보 등)에 데이터를 연계하여 사람을 식별하는 것을 의미함
 - 정보 간접 식별은 나이, 위치 및 고유한 개인 특성과 같이 사람을 식별하는 데 사용할 수 있는 속성을 통해 식별하는 것을 의미함
- 그러나 비식별화는 재식별의 위험을 거의 줄이지 못한다는 지적이 제기됨
 - 비식별화가 이루어진 정보라도 개인 데이터 또는 공개적으로 이용 가능한 정보와 비식별 데이터를 연계할 수 있는 경우 재식별이 발생할 수 있음
 - 1997년 미국 매사추세츠 주의 윌리엄 웰드(William Weld) 주지사의 사례를 보면 누구나 확인할 수 있는 투표인 등록 데이터와 건강유지기구(Health Maintenance Organization)에 등록된 정보를 종합하자 주지사 웰드를 식별할 수 있는 데이터가 모두 파악되었음

| 그림 1 | 주지사 웰드의 정보를 통한 식별화



- 따라서, 직간접적으로 식별되는 데이터를 제거하거나 변환하는 익명화(Anonymization)가 보안 수준을 높이는 데는 효과적임
 - 익명화는 개인의 식별성을 거의 제거하므로 추가적으로 식별성을 제거하기 위한 조치가 필요하지 않으나, 데이터의 유용성을 감소시킨다는 단점이 있음
- 비식별화와 익명화의 장단점을 보완하기 위한 차등 프라이버시(Differential Privacy)* 기술이

등장함

- 차등 프라이버시는 애플(Apple)이 사용자 데이터를 익명화하는 데 이용하겠다고 발표한 이후 인기가 증가했음
- 차등 프라이버시는 데이터 세트에 추가되는 정보의 식별성을 조절할 수 있으며, 빅데이터 분석 과정에서 무작위성을 추가하여 편향도를 줄일 수 있음

*차등 프라이버시: 하버드대학교 신시아 드워크가 2006년에 처음으로 제안한 개인정보를 보호하기 위한 기술로 데이터 통계처리 시 통계값이 바뀌지 않도록 하면서 데이터에 노이즈(Noise)를 섞어 역추적을 하더라도 특정 개인을 식별할 수 없도록 함

- 비식별화, 가명화 및 익명화를 위한 서비스 시장이 성장하고 있다는 점도 주목할 만함
 - 독일 기업 KIProtect의 경우 대규모 데이터 세트를 이용하는 기업이 데이터를 안전하게 보호하고 API를 통해 클라이언트 기업의 데이터 처리와 통합하여 개인 데이터나 민감한 데이터를 탐지하고 보호할 수 있도록 지원함
 - 데이터를 처리하는 기업이 개인 데이터보호를 위해 이러한 서비스 기업에 개인 데이터보호 업무를 아웃소싱할 수 있으며, 이에 따라 자체 사내 개인정보보호 시스템을 구축해야 하는 부담을 줄일 수 있음
- 이러한 개인정보보호 강화 기술이 지속적으로 성장한다면 빅데이터 및 머신러닝 모델의 취약성을 보완하고 효율성과 보안 모두 보장할 수 있을 것이라 기대됨



사후 데이터 관리 책임

- ▶ 자동화된 결정 시스템이 사람의 개입 없이 작동하는 경우 알고리즘과 시스템의 생성, 제조, 운영, 유지 관리 담당자 및 사용자가 프로세스에서 각각의 요소에 대해 책임을 지도록 보장할 필요가 있음
 - 소비자들이 자신의 개인정보를 공유한 뒤 정보나 정보가 정확하지 않게 공유되었거나 개인정보 공유로 인해 불이익을 얻게 되었을 경우, 사후적으로 개입할 수 있는 권리가 필요함
 - 이를 위해서는 투명성이나 추적가능성이 요구되며, 자동화된 결정의 결정권자는 그 결정에 대한 자세한 사유를 설명해야 함
 - 또한, 데이터 처리 및 알고리즘 설계에 대한 각 결정 사항에 대한 문서화가 진행되어야 함
 - 마지막으로 알고리즘과 시스템의 생성, 설계, 제조, 운영, 유지 담당자 및 사용자는 적절한 경우 법적 책임의 형태로 결정에 대한 경제적 책임을 져야 함
- ▶ 개인정보보호 법률에서 보장하는 핵심 보호 권리는 소비자가 본인 데이터에 접근하여 오류를 수정하거나 이의를 제기할 수 있는 권리임
 - 예를 들어, 최근 제정된 2018년 캘리포니아 소비자 개인정보보호법(California Consumer Privacy Act of 2018)은 소비자가 요청할 경우 캘리포니아 거주자의 개인정보를 수집하는 사업자가 전년도 동안 해당 소비자에 대해 수집한 개인정보를 공개하도록 하고 있는데, 여기에는 수집된 정보의 상세 내용과 정보를 공유한 제3자 정보까지 포함됨
 - 일부 국가에서는 개인이 제공된 데이터와 관찰된 데이터 뿐만 아니라 추론된 데이터와 파생 데이터에까지 접근할 수 있는 권리를 보장하고 있음
 - 여기에는 데이터 관리자가 개발한 프로파일, 데이터 처리의 목적, 보유 데이터의 범주 및 출처에 대한 정보가 포함될 수 있음
 - 생년월일, 주소, 급여 수준 또는 혼인 상태 등 간단한 데이터는 이미 수정할 권리가 보장되고 있으며, 추론을 통한 예상 소득 수준, 사망 연령 등 추론 데이터 또한 자동화된 결정에서 중요하므로 일각에서는 추론 데이터까지 수정할 권리가 보장되어야 한다고 주장하고 있음
 - 개인정보 수집 및 처리 데이터 중 추후 필요하지 않은 데이터에 대해서는 개인에게 이를 삭제할 수 있는 권리를 부여하는 국가들이 증가하고 있음
 - EU의 일반데이터보호규정에서는 수집 또는 처리된 목적과 관련하여 데이터가 더 이상 필요하지 않은 경우, 개인이 동의를 철회하고 처리에 대한 다른 법적 근거가 없는 경우 개인의 개인 데이터 삭제 요청 권리를 보장하고 있음
 - 또한 스페인에서는 구글(Google)에 잊혀질 권리를 보장하도록 요구한 사건이 유명함

- 머신러닝을 통해 도출된 추론에 대한 접근권, 수정 또는 삭제 권리를 보장해야 하는지 여부는 아직 확립되지 않았으며, 많은 국가에서도 아직 보장해주지 않음
 - 대부분 국가에서 데이터보호법이 적용되더라도 소비자의 데이터보호보다 머신러닝 처리를 통해 얻을 수 있는 기업의 이익에 더 큰 비중을 둘 것으로 보임
- ▶ 빅데이터와 머신러닝을 활용한 자동화된 결정에 문제가 있을 경우 소비자의 알권리 상 설명이 필요한데, 기술적 그리고 법적 난관이 있어 관심을 두고 개선해야 함
- 소비자에게 자동화된 결정에 대한 설명을 해주는 과정에서 두 가지 문제가 발생함
 - 첫 번째로, 머신러닝은 컴퓨터 프로그래머가 일일이 지시해서 결과를 도출하지 않기 때문에 머신러닝으로 도출된 결과에 대해 소비자에게 쉽게 설명하기가 매우 어려움
 - 두 번째로, 머신러닝 모델은 기업에서 투자를 통해서 마련한 자산이기 때문에 영업비밀과 소프트웨어 저작권 문제로 모든 사항을 공개하기 어려움
- 그럼에도 불구하고 소비자의 알 권리를 보장해주기 위한 노력이 이루어져야 함
 - 사회의 각종 서비스가 자동화된 의사결정 또는 빅데이터에 의존하고 있으므로 사회 전반에서 합리적으로 결정이 내려지고 있다는 신뢰가 있어야 함
 - 실제로 일반인들은 머신러닝 모델을 이해하지 못할 뿐만 아니라 머신러닝을 개발하는 개발자조차 머신러닝 알고리즘에 대해 설명하지 못하는 경우가 많음
 - 따라서, 의학, 은행, 보험 및 기타 분야 종사자들은 머신러닝 모델을 이해하고 운영 방식에 대한 투명성을 제고해야 함
- 국가적 차원에서도 자동화된 결정에 대한 관심을 기울이고 있음
 - 소비자는 자동화된 결정에 사용되는 기준과 절차에 대한 명확하고 관련 있는 정보를 요청할 권리를 보장받아야 함
 - 브라질의 2018년 데이터보호법(Data Protection Act 2018)은 소비자의 이익에 영향을 미치는 개인 데이터의 자동 처리를 기반으로 한 결정의 검토를 요청할 권리를 제공함
 - 브라질 뿐만 아니라 다른 국가에서도 개인정보보호법에 따라 자동화된 결정에 대해 더 엄격한 조사를 요구하고 있음
 - EU의 경우 데이터 관리자가 자동화된 결정과 관련된 논리에 대해 소비자에게 의미 있는 정보를 제공해야 한다고 권고함
- 일각에서는 개인정보보호법이 실질적으로 소비자에게 ‘합리적 추론에 대한 권리(Right to reasonable inferences)’를 보장해야 한다고 주장함
 - 이러한 권리가 부여된다면 머신러닝을 통한 추론이 소비자에게 불리한 결정 또는 평판 손상으로 이어지고, 프라이버시를 침해할 리스크가 높은 경우 데이터 관리자는 자동화된 결정 및 처리에 따른 추론 관련성(the relevance of the inferences)에 대해 사전 설명해야 함

- 의사결정이 이루어진 후 소비자는 이의를 제기할 권리가 주어짐
- 전반적으로, 소비자들에게 개인정보 침해와 빅데이터 및 머신러닝의 위험에 대한 효과적인 구제책을 제공하기 위해 국가에서 수행해야 할 일이 많이 남아있음

PART VI

리스크 관리와 데이터 윤리 기반 프레임워크 구축

- ▶ 개인정보침해 및 데이터 보안 사고 위험을 완화하기 위해 데이터를 관리하는 기업에서는 리스크 관리 프레임워크와 프로세스를 적용해야 함
 - 이러한 프레임워크와 프로세스는 다른 리스크와 마찬가지로 소비자 프라이버시 및 차별과 관련된 리스크를 평가하기 위해 사용될 수 있음
 - 미국 국립표준기술연구소(NIST, The US National Institute of Standards and Technology)는 2020년 1월 리스크 관리 접근법에 초점을 맞춘 프레임워크 연구 버전 1.0 보고서를 발표하여, 기업들이 기술의 컨셉을 잡을 때 개인정보보호를 고려한 관행이 이루어질 수 있도록 기업 리스크관리를 통한 개인정보보호 개선 방법에 대해서 다룸
 - 이 프레임워크는 즉각적인 규정 준수보다 리스크 관리의 우선 순위를 지정하는 것의 중요성을 강조함
 - 머신러닝 시스템의 리스크 관리 프로세스에는 목표를 문서화하고, 머신러닝 모델의 개발 및 시험과 검증 및 법적 검토, 수명 전반에 걸쳐 모델에 대한 주기적인 감사 등 리스크 관리를 위한 여러 사항이 포함됨
 - 이러한 리스크 관리에는 프로세스 관리뿐만 아니라 입력 및 출력 데이터, 알고리즘 생성 및 운영 관리도 포함되어야 함
 - 입력 데이터 측면에서 리스크 관리를 위해서는 주변 시스템의 데이터 의존도, 개인 데이터 활용 목적 및 방법, 개인 데이터보호 방법(암호화 등)을 문서화해야 하며, 입력 데이터 자체의 완전성, 정확성, 일관성, 가용성(completeness, accuracy, consistency, availability) 등을 확인해야 함
 - 출력 데이터 측면에서 리스크 관리를 위해서는 머신러닝 모델이 편향적인 결과를 도출해내는 것을 막기 위해 출력 데이터 모니터링과 차별 방지 메커니즘 도입 등이 필요함
 - 전체적으로 데이터 관리자는 머신러닝 모델이 데이터 및 개인정보보호법, 차별 방지법을 위반할 가능성이 있는 경우 머신러닝 모델을 수정하거나 사용을 중단해야 함
- ▶ 빅데이터 및 머신러닝에서 개인정보보호를 효과적으로 수행하기 위해서 기업은 법률과 규정에 앞서 사생활 침해 최소화할 수 있는 제품과 서비스를 고안하려는 원칙을 갖춰야 함
 - 사생활 침해 최소화를 위해 앤 카부키안(Ann Cavoukian)은 프라이버시 7가지 원칙을 제시하였음
 - ① 사생활 침해 발생이 이루어지기 전에 사전 예방적인 조치를 취해야 함
 - ② 소비자가 프라이버시를 보호하기 위해 설정을 변경할 필요가 없도록 프라이버시를 기본 설정으로 함

- ③ 시스템 고안 과정에서 프라이버시를 주요 사항으로 고려함. 프라이버시는 시스템 설계 후 조정되어서는 안되며, 기능성을 약화시켜서도 안됨
- ④ 고객 신뢰 강화, 데이터 침해로 인한 위험 감소 등의 이점을 누리는 윈-윈 접근 방식을 채택함
- ⑤ 엔드 투 엔드(End-to-End) 보안을 통해 데이터의 수명 주기에 걸쳐 안전한 수집, 저장 및 파기를 보장함
- ⑥ 내부 모니터링 및 독립적인 검증이 가능하도록 정책을 설정하고 기록을 보관하며, 가시성과 투명성을 보장함
- ⑦ 소비자에게 정보에 대한 접근 권한과 데이터의 수정 및 업데이트를 요구할 수 있는 권리를 보장함

▶ **첨단 기술 산업 전반에서는 인공지능과 머신러닝에 대한 윤리 의식을 함양하기 위한 노력이 이루어지고 있음**

- 컴퓨팅 기계협회(ACM, The Association for Computing Machinery)와 전기전자학회, 인공지능 파트너십(Partnership on AI), 소프트웨어 정보 산업협회(SIIA, Software & Information Industry Association) 등 기관뿐만 아니라 구글과 마이크로소프트(Microsoft) 등 우수 기업에서는 개인정보보호와 지속가능한 성장, 데이터의 투명성 제고를 위해 노력하고 있음
- 특히 금융 서비스 분야에서는 스마트 캠페인(Smart Campaign)을 통해 디지털 신용 표준의 일부로 알고리즘 및 데이터 기반 자동 의사 결정 지표 초안을 마련했으며, 금융 서비스 소비자 데이터보호를 위한 자체 규제 표준을 개발하였음
 - 고객보호인증을 수행하는 독립 평가 기관인 MFR은 스마트 캠페인 참여를 통해 디지털 신용제공자를 위한 고객 보호 기준을 마련함
 - MFR은 케냐의 금융 서비스 제공업체인 4G Capital과 탈라(Tala)와 고객 보호 기준을 시험 운용한 뒤 결과를 반영해 개정된 표준을 발표함
- 이러한 조치만으로는 공정성, 책임성 및 투명성을 확보하기 어렵지만, 리스크 관리를 위한 신속한 커뮤니케이션을 가능하게 하며 소비자 개인정보보호 조치의 발전에 기여한다는 점에서 의의가 있음

PART VII

결론

- ▶ 본 보고서는 빅데이터 및 머신러닝 기술과 개인정보보호 법률 및 규제가 직면한 다양한 과제를 탐구함
 - 특히 이러한 기술은 기업이 소비자에게 제공하는 서비스에서 사용된다는 점에서 매우 중요하며, 부정확한 데이터 입력 및 머신러닝을 기반으로 한 추론의 편향(Biased Inference)과 차별적 취급에 따른 리스크도 존재함
 - 그러나 알고리즘을 통한 추론과 자동화된 결정(Automated Decision)에 대한 투명성을 보장하거나 인공지능으로 인해 발생한 소비자 피해를 입증하는 것은 매우 어려움
 - 해당 사항은 현재 소비자 데이터보호, 개인정보보호 법률 및 규정에서 대부분 다루고 있지 않으며, 향후 빅데이터 및 머신러닝 기술의 활용이 증가할 것으로 전망되기 때문에 규제 기관에서 관심을 가져야 할 것으로 보임
- ▶ 빅데이터와 머신러닝 전반에 적용할 수 있는 표준 및 기술을 개발하고, 데이터 프라이버시 보호와 혁신 간의 균형을 맞추기 위해 아래와 같은 여러 노력을 기울여야 함
 - 개인정보의 이용 및 공유에 대해 소비자에게 알기 쉽게 설명을 제공하는 등 소비자 공지 방식 개선, 소비자 개인정보보호를 강화하는 데 필요한 조치 마련
 - 무분별한 소비자 데이터 공유를 막기 위해 개인정보 사용 및 공유 관련 규제 강화
 - 인공지능 및 머신러닝 모델 설계에 프라이버시 원칙을 적용하기 위한 표준 개발
 - 사전 예방적 설계 접근, 프라이버시 기본 설정 적용, 설계에 의한 프라이버시 채택, 소비자 신뢰 지향, 정보에 대한 종단 간 보안, 소비자 접근성 개선 등 표준 개발이 필요함
 - 규제 기관에서 검토가 가능하도록 설계 프로세스가 생성, 기록 및 보고되어야 하며, 데이터에 대한 확인도 수반되어야 함
 - 인공지능 컴퓨터 프로그래밍에 대한 윤리적 표준 개발, 신규 문제 식별 및 접근 방식 개발
 - 추론의 신뢰도 확립을 위한 표준 개발
 - 출력 데이터의 평가와 머신러닝 모델의 개인정보보호 및 차별 방지 원칙에 대한 준수 여부 확인을 통해 추론의 신뢰도를 보장할 수 있음
 - 머신러닝 시스템이 도출한 추론에 사용된 데이터 정보 및 추론 프로세스를 설명하기 위한 표준 개발

- 자동 결정에 대한 소비자 이의 제기 절차 마련 및 인공지능 및 머신러닝 모델의 설계 및 운영에 대한 책임 평가를 위한 프레임워크 마련
- ▶ 우리나라의 경우에도 정부 차원으로 빅데이터와 머신러닝 등에서의 데이터 보호기술 발전을 추진하고 있으며 향후로도 해당 분야에 대한 투자는 계속될 것으로 예측됨
- 개인정보보호위원회가 발표한 개인정보보호·활용 기술 R&D 로드맵에서는 데이터보호를 위한 주요 기술과 세부 추진 방안을 제시함⁸⁾

기술	세부 추진 방향
개인정보 동의 관리 기술	개인정보 수집, 제3자 제공에 대한 동의, 동의 철회 등 정보 주체의 자기정보 통제권 보장을 위한 UX 중심의 설계 방안 마련
정보 주체의 온라인 활동 기록 통제 기술	온라인 상의 개인 행태정보 수집을 방지하고 개인 행태정보의 추적을 승인 또는 중단할 수 있도록 지원
다크웹에서의 개인정보 불법거래 추적 및 차단 기술	다크웹 기반의 기존 불법거래 사례 분석을 통해 다크웹 내 개인정보 불법거래 추적기술 개발
비정형 데이터에서 개인정보 탐지 기술	대화형 텍스트, 음성 및 영상 데이터 등 비정형 데이터에 대한 개인정보 식별을 위해 필요한 기술 개발
개인정보 파편화 및 결합 기술	이미지·영상 정보에 포함된 개인정보를 실시간으로 파편화하여 저장하고, 필요시 이를 실시간으로 결합하여 개인정보를 복구·관리할 수 있는 기술 개발
차세대 가명·익명처리 및 결합 기술	금융, 쇼핑 등 거래 기반의 실시간 트랜잭션 데이터 등 다양한 형태의 데이터에 적합한 가명·익명 처리 기술 및 통계, 암호, 연산 기법과 연계한 가명·익명처리 기술 개발
가명·익명정보 안전성 평가 기술	데이터 유형별 가명·익명처리 기술의 재식별 가능성 및 공격모델 등을 연구하여, 기술의 안전성 검증
개인정보 변조 및 재현데이터 생성 기술	인공지능 기술 선점 및 개인정보 침해 최소화를 위해 텍스트, 음성, 영상 등 비정형 데이터에 대한 유형별 재현데이터 생성 기술 개발
마이데이터 처리 및 관리 기술	개인정보보호법 개정안의 개인정보 전송요구권을 기반으로 마이데이터가 구현될 수 있도록 기반을 제공하는 기술 개발

8) 개인정보보호위원회, 개인정보보호·활용기술 R&D 로드맵, 2021.11

참 고 문 헌

- Financial Inclusion Global Initiative, Big data, machine learning, consumer protection and privacy, ITU, 2021.04
- GPFI, Use of Alternative Data to Enhance Credit Reporting to Enable Access to Digital Financial Services by Individuals and SMEs operating in the Informal Economy, Guidance Note, 2018
- 권정수, “개인정보보호·활용 기술 R&D 로드맵 공개…11대 핵심기술 개발 추진”, IT Daily, 2021년 11월 10일, <http://www.itdaily.kr/news/articleView.html?idxno=205029>
- 김가은, “'디지털 플랫폼 정부'는 AI가 지킨다…최신 보안기법 도입으로 '新보안체계' 구축”, TechM, 2022년 7월 14일, <https://www.techm.kr/news/articleView.html?idxno=99371>
- 개인정보보호위원회, 개인정보보호·활용기술 R&D 로드맵, 2021.11



데이터산업 동향 이슈 브리프

| 발행일 2022년 10월 31일

| 발행처 **K data** 한국데이터산업진흥원

서울시 중구 세종대로 9길 42, 부영빌딩 8층

| 기획 및 편집 데이터산업본부 산업기획팀

| 문의처 Tel: 02-3708-5363, 5364

ISSUE BRIEF

* 본 지에 실린 내용은 한국데이터산업진흥원의 공식 의견과 다를 수 있습니다.
본 내용은 무단전재를 금하여, 가공/인용할 경우 반드시 출처를 밝혀주시기 바랍니다.