

인공지능확산의 국가별 비교연구: 기술혁신, 정책체계, 사회적수용의 삼각구조 분석

이 선 형*

요약

인공지능 기술의 글로벌 확산 양상과 한국 사회의 수용 특성을 규명하기 위해, 본 연구는 기술혁신-정책체계-사회적 수용(T-P-S)의 삼각구조 분석틀을 개발하고 미국·EU·중국·한국의 확산 전략을 비교하였다. Rogers의 혁신확산이론과 Geels의 사회기술체계 이론을 통합하여 2019~2024년 발간된 35편의 정책 문서·산업보고서·학술문헌을 PRISMA 절차에 따라 구조화된 내러티브 리뷰로 분석하였다. 분석 결과, 미국은 혁신주도형, EU는 규제중심형, 중국은 정부주도형, 한국은 혼합형 모델을 추구하는 것으로 나타났다. 정책프레임워크가 기술혁신과 사회적 수용을 매개하며 확산 경로를 차별화하는 핵심 요인으로 확인되었다. 한국은 높은 디지털 인프라와 기술 수용성을 보유하나 기술혁신역량은 중위권에 머물며, 원천기술 확보와 글로벌 표준 선점이 과제로 제기되었다. 본 연구는 T-P-S 통합 분석틀을 제시하고 국내 정보화정책 연구를 종합함으로써 AI 거버넌스 정책 수립과 국가 간 비교연구에 이론적·실증적 기여를 제공한다.

주제어 : 인공지능, AI 확산, 정책체계, 기술혁신, 사회적수용, AI 리터러시

A Comparative Study of AI Diffusion Across Nations: Analyzing the Triangular Structure of Technological Innovation, Policy Framework, and Social Acceptance

Lee, Sun-Hyung*

Abstract

How is AI technology diffusing globally, and how is Korean society accepting it? This study develops an integrated Technology-Policy-Society (T-P-S) analytical framework by combining Rogers' innovation diffusion theory and Geels' socio-technical system theory, and compares AI diffusion strategies across the U.S., EU, China, and South Korea. Following PRISMA guidelines, we conducted a structured narrative review of 35 policy documents, industry reports, and academic literature published from 2019 to 2024. Results reveal four distinct national models: the U.S. pursues an innovation-driven approach, the EU adopts a regulation-centered model, China follows a government-led strategy, and South Korea implements a hybrid model. Policy frameworks serve as key mediators between technological innovation and social acceptance, differentiating diffusion pathways across nations. South Korea demonstrates strong digital infrastructure and high technology acceptance but ranks mid-tier in innovation capacity, facing challenges in securing foundational technologies and leading global standards. This research contributes to AI governance policy development and cross-national comparative studies by presenting the T-P-S integrated framework and synthesizing domestic informatics policy research.

Keywords : artificial intelligence, AI diffusion, policy framework, technological innovation, social acceptance, AI literacy

I. 서론

1. 연구배경

인공지능(Artificial Intelligence; AI)의 역사는 이미 70년을 넘어섰다. 그러나 2020년대 들어 ChatGPT, Claude, Gemini와 같은 대화형 AI가 등장하면서 상황이 완전히 달라졌다. 프로그래밍 지식 없이도 자연어로 AI와 대화하며 다양한 작업을 수행할 수 있게 되었고, 이는 AI 확산의 결정적 전환점이 되었다(Brown et al., 2020; Radford et al., 2019).

각국 정부도 본격적으로 대응에 나섰다. 미국은 2023년 10월 「안전하고 신뢰할 수 있는 AI 개발과 사용에 관한 행정명령」을 발표했고(White House, 2023), EU는 2024년 6월 세계 최초로 포괄적 AI 규제법인 「AI Act」를 확정했다(European Commission, 2024). 중국은 이미 2017년부터 「차세대 인공지능 발전 계획」을 통해 체계적인 국가 전략을 추진해왔으며(State Council of China, 2017), 한국도 2019년 「AI 국가전략」과 2021년 「신뢰할 수 있는 인공지능 윤리기준」을 마련했다(관계부처합동, 2019; 과학기술정보통신부, 2021).

한국은 미·중 기술 경쟁이 심화되는 환경 속에서 독자적인 AI 생태계 구축과 기술 경쟁력 확보를 모색하고 있다. 특히 반도체와 디지털 인프라 분야의 강점을 AI 산업 발전과 연계하기 위한 정책적 노력이 지속되고 있다.

흥미로운 점은 동일한 기술임에도 국가마다 채택 속도와 활용 방식이 뚜렷하게 다르다는 점이다. 기술 자체의 우수성만으로는 이런 차이를 설명하기 어렵다. 각국의 제도적 맥락과 사회문화적 배경은 AI 확산 경로에 중요한 영향을 미치는 것으로 지적된다(Cath, 2018). 한국은 우수한 디지털 인프라를 바탕으로 AI 선도국 유형에 포함되지만, AI 인적자본과 기술혁신역량은 중위권 수준에 머무르고 있어 장기적 역량 강화가 요구된다(신승윤, 2024).

2. 문제제기 및 연구 필요성

기존 AI 연구는 대체로 기술 성능 향상이나 특정 정책 효과, 기업 도입 사례 등 미시적 수준에 초점을 두어 왔다. 기술 수용 모델(TAM)(Davis, 1989)이나 혁신 확산 이론(IDT)(Rogers, 2003) 같은 전통적 이론들은 개인이나 조직 수준 분석에는 유용하지만, AI의 블랙박스 특성, 급속한 진화, 광범위한 사회적 파급효과를 종합적으로 설명하기에는 한계가 있다.

국가 간 AI 확산 양상을 체계적으로 비교한 연구도 부족하다. OECD(2023)나 Cath(2018)의 연구가 국가별 AI 전략을 다루고 있지만, 대부분 정책 내용 비교에 그친다. 본 연구는 기술-정책-사회 상호작용을 통한 확산 경로의 차별화라는 통합적 관점을 제시하려 한다. AI 확산을 국가 간 비교의 관점에서 기술·정책·사회 요소를 통합해 설명하려는 시도는 아직 부족하며, 이는 정책 설계와 국제 경쟁 전략 수립에 중요한 공백으로 지적된다.

국내에서는 최근 AI 윤리와 사회적 수용성에 대한 관심이 높아졌다. 최근 AI의 확산과 함께 윤리, 책임성 및 사회적 신뢰에 대한 논의가 확대되고 있으며, AI 정책은 기술적 성과뿐 아니라 사회적 수용성과 제도적 정당성을 함께 고려하는 방향으로 발전하고 있다. 그러나 이런 규범적 논의가 실제 정책 집행과 산업 현장에서 어떻게 구현되는지, 국제 비교 관점에서 한국 모델의 특수성이 무엇인지에 대한 실증 분석은 여전히 부족하다.

한국 사례는 특히 주목할 만하다. 한국지능정보사회진흥원(2024)에 따르면 2023년 한국인의 생성형 AI 서비스 경험률은 17.6%였으며, 특히 20대(33.7%)와 30대(28.2%)에서 높은 활용도를 보였다. 문제는 이런 빠른 확산 속도가 질적 성숙도나 사회적 준비도와 일치하는지, 장기적으로 어떤 정책적 함의를 갖는지 명확하지 않다는 점이다.

3. 연구목적 및 접근방법

본 연구는 세 가지 목적을 갖는다. 첫째, AI 확산을 기

술 혁신-정책 체계-사회적 수용의 삼각 구조로 분석하는 통합 프레임워크를 만든다. 둘째, 이를 바탕으로 미국, EU, 중국, 한국의 AI 확산 전략과 성과를 비교한다. 셋째, 한국의 AI 확산 특성을 규명하고 정책적 시사점을 도출한다.

이러한 문제의식을 바탕으로 다음과 같은 연구문제를 설정하였다.

- RQ1. AI 확산을 설명하는 통합적 분석틀은 어떻게 구성될 수 있는가?
 RQ2. 미국, EU, 중국, 한국의 AI 확산 전략과 성과는 어떤 차이를 보이는가?
 RQ3. 한국의 AI 확산 경로는 어떤 특성을 지니며, 어떤 정책적 시사점을 줄 수 있는가?

본 연구는 Rogers(2003)의 혁신 확산 이론과 Geels(2002)의 사회기술체계 이론을 AI 맥락에 적용하여 분석틀을 구성하였다. Rogers 이론은 기술의 개별 특성과 초기 확산 과정을 설명하는 데 쓰고, Geels의 다층적 관점(Multi-Level Perspective)은 기술 니치, 사회기술 레짐, 거대 경관의 상호작용을 통해 기술-정책-사회 요소 간 거시적 시스템 변화를 분석하는 데 활용했다.

분석 자료는 2019~2024년 발간된 정책 문서, 산업 보고서, 학술 논문이다. 이 시기를 선택한 이유는 명확하다. 2019년은 한국 「AI 국가전략」이 나온 해로 본격적인 정책 대응의 시작점이다. 2020년은 GPT-3 등장으로 대규모 언어 모델 시대가 열린 해다. 2022~2023년은 ChatGPT로 대표되는 생성형 AI가 대중화되며 확산이 급격히 가속화된 전환기다. 2024년은 주요국의 규제 법안이 구체화되는 제도적 성숙기에 해당한다.

II. 이론적 배경

본 장에서는 인공지능 기술의 발전 경로와 확산 메커니즘, 정책·거버넌스 논의, 그리고 사회적 수용 및 국가 비교 연구를 검토한다. 이를 통해 본 연구의 분석틀

인 기술혁신(Technology)-정책체계(Policy)-사회적수용(Society) 삼각구조의 이론적 토대를 마련하고자 한다.

1. AI 기술 발전의 단계적 전환

AI 발전사는 몇 차례 중요한 패러다임 전환을 거쳤다. 1950년대부터 1980년대까지는 인간 전문가의 지식을 규칙으로 코딩하는 방식이 주를 이뤘다. MYCIN이나 DENDRAL 같은 전문가 시스템(Buchanan & Shortliffe, 1984)은 제한된 영역에서 성과를 보였으나, 데이터 학습 없이 규칙 기반으로만 작동하여 유연성이 부족했다. 1980년대 후반 'AI 겨울'이라 불리는 침체기가 이어졌다.

1990년대 이후 데이터를 통해 패턴을 학습하는 기계학습(machine learning)이 부상했다. 다만 어떤 특징(feature)을 추출할지 여전히 인간이 결정해야 했고, 도메인 지식 의존도가 높아 범용성에 제약이 있었다. 2010년대 초반 심층학습(deep learning)이 등장하면서 원시 데이터로부터 계층적 특징을 자동 학습하는 방식이 본격화되었고, Krizhevsky et al.(2012)의 Alex-Net을 계기로 이미지 인식·음성 인식·자연어 처리 등에서 성능 도약이 관찰되었다. Vaswani et al.(2017)이 제안한 트랜스포머(Transformer) 아키텍처는 병렬 처리를 통해 대규모 데이터 학습을 가능하게 만들어 현재의 대규모 언어 모델(LLM)로 이어졌으며, Brown et al.(2020)의 GPT-3는 이러한 발전의 대표적 사례로 언급된다.

이러한 기술적 진화는 단순한 성능 향상을 넘어 AI의 사회적 위상을 변화시켰다. 전문가의 도구에서 일반 대중이 접근 가능한 서비스로, 특정 과제 해결 도구에서 범용 지능을 지향하는 시스템으로, 그리고 보조 도구에서 자율적 행위자(autonomous agent)로의 전환이 진행 중이다. 이러한 변화는 기술 확산의 속도와 범위를 결정하는 주요 요인이며, 동시에 새로운 정책적 과제를 제기한다(Dwivedi et al., 2021).

2. 기술 확산의 이론적 관점

Rogers(2003)의 혁신확산이론은 상대적 이점, 적합성, 복잡성, 시험가능성, 관찰가능성이라는 다섯 가지 특성에 주목하며, 혁신 수용 집단이 시간에 따라 S자형 곡선을 그리며 확산된다고 본다. AI 확산에 이 관점을 적용하면, 생성형 AI는 뚜렷한 상대적 이점과 높은 시험가능성을 갖는 반면, 블랙박스 구조와 윤리·법적 우려는 관찰가능성과 적합성 인식을 제약한다.

Davis(1989)의 기술수용모델(TAM)은 인지된 유용성과 인지된 사용 용이성을 기술 수용의 핵심 요인으로 설정한다. 이 모델은 개인 수준의 수용태도와 행동의도 분석에 강점을 가지지만, 국가 정책, 제도 신뢰, 사회문화적 가치 등 거시적 맥락을 충분히 반영하지 못한다는 비판을 받아왔다.

AI 윤리와 거버넌스 논의는 기술적 안전성뿐 아니라 사회적 신뢰와 책임성을 함께 확보해야 한다는 방향으로 전개되고 있다. 국가 간 기술 경쟁과 디지털 전환 환경의 변화는 각국의 AI 정책 방향과 제도 설계에 영향을 미치는 중요한 외부 요인으로 논의되고 있다.

본 연구는 Rogers의 다섯 속성을 기술혁신, 정책체계, 사회적 수용의 세 축과 연계하여 조작화한다.

〈표 1〉의 조작화 구조를 구체적으로 살펴보면, 상대적 이점은 기술 차원에서는 알고리즘 성능으로, 정책 차원에서는 경제적 효과 측정으로, 사회 차원에서는 편의성 체감으로 구체화된다. 시험가능성의 경우, 기술 차원의 파일럿 테스트가 정책 차원의 규제 샌드박스를 거쳐 사회적 차원의 단계적 도입 기회로 연결되는 구조이다.

이러한 T-P-S 조작화는 Rogers(2003)와 Geels(2002)의 통합에 기반한다. Rogers(2003)의 혁신확산 이론은 개인과 조직 수준의 기술 수용 과정을 설명하나 제도·규범 등 구조적 맥락을 충분히 다루지 못한다. 반면 Geels(2002)의 사회기술체계 이론은 니치(niche), 레짐(regime), 경관(landscape)의 다층적 상호작용을 통해 거시적 전환을 설명하나 미시적 의사결정 동학을 간과한다. AI 확산 맥락에서 보면, 생성형 AI는 니치 단계의 실험을 빠르게 통과해 기존 레짐을 재구성하는 압력을 가하고, 디지털 전환·팬데믹·지정학적 경쟁과 같은 거시적 경관 변화가 이를 가속화한다.

본 연구는 Rogers의 5가지 혁신 속성을 T-P-S 구조로 조작화하고, Geels의 다층적 관점을 정책체계(P)가 기술혁신(T)과 사회적 수용(S)을 매개하는 레짐으로 설정함으로써 미시-거시 연결고리를 제공한다.

〈표 1〉 Rogers 혁신 속성의 T-P-S 연계 조작화

Rogers 혁신 속성	정의	T 기술혁신	P 정책체계	S 사회적 수용
상대적 이점 (Relative Advantage)	기존 대안 대비 우수성 인식	알고리즘 성능, 효율성, 처리 속도	경제적 효과 측정, ROI 가시성	편의성, 생산성 체감, 실용적 가치
적합성 (Compatibility)	기존 가치·경험·필요와의 부합성	기존 시스템·인프라와의 통합 가능성	법적 규범·제도와의 정합성	문화적 가치·윤리 기준과의 조화
복잡성 (Complexity)	이해·사용의 난이도	기술 구조의 복잡도, 블랙박스 특성	규제 요구사항의 복잡도	학습 곡선, 사용자 인터페이스 난이도
시험가능성 (Trialability)	제한적 범위에서 실험 가능성	파일럿 테스트, PoC (Proof of Concept)	규제 샌드박스, 시범사업	무료 체험판, 단계적 도입 기회
관찰가능성 (Observability)	결과의 가시성 및 전달 가능성	성과 지표, 벤치마크 측정	정책 평가, 모니터링 체계	타인 경험 관찰, 신뢰 형성

출처: Rogers(2003)의 혁신확산이론을 본 연구의 T-P-S 프레임워크에 맞게 재구성.

이러한 조작용은 동일한 기술적 특성이 정책 환경과 사회적 맥락에 따라 상이한 방식으로 확산에 영향을 미친다는 본 연구의 핵심 논지를 뒷받침한다.

3. 정책과 거버넌스의 역할

AI 확산에서 정책과 거버넌스의 역할은 규제를 넘어 생태계 조성, 신뢰 구축, 국제 협력에 이른다. Cath (2018)는 AI 거버넌스를 기술적, 법적, 윤리적 세 차원으로 구분하며, 각 차원이 상호 보완적으로 작동해야 효과적인 정책 집행이 가능하다고 주장했다.

OECD(2023)는 회원국의 AI 원칙 이행 현황을 평가하며, 대부분의 국가가 원칙 수립 단계를 넘어 구체적인 집행 메커니즘을 마련하는 단계로 진입했다고 분석했다. 국가 간 정책 접근의 차이는 여전히 크다. 김병운(2016)은 한국의 AI 거버넌스·R&D·법·제도 환경을 종합적으로 진단하며, 국가 전략이 기술 생태계 형성과 확산의 구조적 전제조건이라고 강조하였다. 이는 정책 체계(P)가 기술혁신(T)과 사회적 수용(S)을 실질적으로 연결하는 매개 역할을 한다는 본 연구의 분석 관점을 지지한다.

미국은 행정명령과 가이드라인 중심의 유연한 접근을 취하는 반면(White House, 2023), EU는 위험 기반 규제를 통한 강제적 준수 체계를 구축했다(European Commission, 2024). 중국은 국가 전략 산업으로 AI를 육성하면서 동시에 콘텐츠 규제를 통한 사회 통제를 강화하는 이중 전략을 추구한다(State Council of China, 2017).

한국의 정책 접근은 이들 모델의 절충형 성격을 띤다. 관계부처 합동(2019)의 「AI 국가전략」은 기술 육성, 산업 진흥, 사회 안전망을 균형 있게 다루며, 과학기술정보통신부(2021)의 윤리 기준은 자율 준수를 원칙으로 하되 필요 시 법제화 가능성을 열어두고 있다. 이러한 접근은 규제와 혁신의 균형을 모색하는 것으로 평가받으나, 명확한 방향성 부재라는 비판도 받는다.

한국 AI 정책의 과제로는 원천 기술 확보를 위한 장

기 R&D 투자 강화, 산학연 협력 생태계의 실질적 작동, 그리고 국제 표준 선점을 위한 전략적 참여가 반복적으로 제기된다. 한국지능정보사회진흥원(2024)의 최신 보고서는 AI 윤리 원칙의 구체적 이행을 위해 산업별 가이드라인 개발, 자율 점검 도구 보급, 인증 제도 도입 등 단계적 접근을 제안하며, 이는 자율 규제와 공적 개입의 조화를 추구하는 한국형 거버넌스 모델의 구체화로 해석될 수 있다.

4. 사회적 수용과 국가 비교 연구

AI 확산의 성패는 기술·정책뿐 아니라 사회적 수용과 디지털 리터러시 수준에 좌우된다. 국내외 실증 연구들은 AI에 대한 기대와 우려가 동시에 존재하며, 특히 공공서비스나 정책 의사결정 영역에서 개인의 위험 인식, 신뢰 수준, 리터러시가 정책 수용 의도에 유의미한 영향을 미친다는 점을 반복적으로 확인하고 있다.

국가 비교 관점에서 AI 정책과 확산 경로를 다룬 선행연구들은 미국, EU, 중국, 한국 등을 중심으로 각국의 전략과 제도적 특성을 정리해 왔다. 다만 많은 연구가 개별 국가의 사례 서술이나 정책 소개 수준에 머물러, 기술·정책·사회 요소를 통합한 구조적 비교는 상대적으로 부족한 편이다. 신승윤(2024)은 OECD 38개국을 AI 인적자본, AI 기반 인프라, 기술혁신역량 요인을 기준으로 유형화하고, 한국이 “AI 선도국” 그룹에 포함되지만 인적자본과 기술혁신역량 측면에서는 중위권에 머무른다는 점을 지적하였다. 일부 연구는 미국을 혁신 주도형, EU를 규제중심형, 중국을 정부주도형, 한국을 혼합형으로 분류하지만, 이러한 유형화가 어떤 이론적 틀에 기반하는지, 각 유형이 사회적 수용과 어떤 방식으로 연결되는지에 대해서는 충분히 모델링되지 않은 경우가 많다.

기존 선행연구들은 AI 기술의 패러다임 전환, 정책·거버넌스의 제도적 발전, 사회적 수용과 리터러시 요인, 그리고 국가 간 정책 비교를 각각 독립적으로 분석해왔다. 정보화정책 분야의 최근 연구들 역시 공공

〈표 2〉 AI 확산 관련 선행연구 요약표

구분	연구자/연도	국가·대상	방법	주요 결과	본 연구의 차별성 및 기여
기술혁신	Buchanan & Shortliffe (1984)	미국—전문가시스템	기술 사례 분석	규칙 기반AI의 구조적 한계 제시	본 연구는 이를 기술 패러다임 전환의 출발점으로 설정하고, 규칙→기계학습→심층학습→LLM으로 이어지는 기술혁신(T)의 역사적 경로를 체계화
	Krizhevsky et al. (2012)	글로벌—이미지 인식	모델 개발	딥러닝 기반AI 성능 도약	선행연구는 기술 성능에 집중. 본 연구는 이를 T축의 핵심 분기점으로 설정하고, 기술혁신이 정책(P)·사회(S)와 상호작용하며 확산 경로를 구조적으로 변화시키는 메커니즘을 분석
	Vaswani et al. (2017)	글로벌—언어 모델	모델 제안	트랜스포머·LLM 기반 제시	선행연구는 모델 아키텍처 제안. 본 연구는 이를 생성형AI 확산의 기술적 분기점으로 해석하고, 레짐 변화(Geels, 2002)와 연계하여 T-P-S 통합 분석의 이론적 근거로 활용
정책·거버넌스	Cath (2018)	글로벌	문헌 분석	기술·법·윤리 삼중 거버넌스 필요	선행연구는 거버넌스 차원 제시. 본 연구는 이를 T-P-S 프레임워크의 P축으로 조작화하고, 정책이 기술혁신(T)과 사회적 수용(S)을 매개하는 구조적 메커니즘임을 실증
	OECD (2023)	OECD 회원국	비교 정책 분석	원칙 중심 → 실행 메커니즘 강화	선행연구는 정책 이행 현황 비교. 본 연구는 동일 데이터를 활용하되 한국의 혼합형 모델을 미국·EU·중국과 T-P-S 기준으로 구조적 비교하여 차별화된 정책 시사점 도출
	신승윤(2024), 정보화정책	OECD 38개국	퍼지셋 이상형 분석	국가별AI 경쟁력 8유형 분류	선행연구는 정태적 유형 분류. 본 연구는 한국이 '선도국이나 중위권'이라는 신승윤의 발견을 T-P-S 동태적 확산 경로와 연계하여, 혼합형 모델의 강점·한계를 정책적으로 해석
	은종환·황성수 (2020), 정보화정책	한국—정책사례	탐색적 분석	문제구조화 수준에 따라AI 도입 성패 결정	선행연구는 공공부문 사례. 본 연구는 이를 P축이 T축을 현장에서 작동시키는 핵심 조건이라는 이론적 논거로 확장하고, 한국 혼합형 모델의 실행 메커니즘 분석에 활용
사회수용·리터러시	장창기·성옥준 (2022), 정보화정책	한국 공공서비스	구조방정식	인식·신뢰·리터러시가 수용 의도 결정	선행연구는 개인 수준 실증. 본 연구는 이를 S축의 실증 기반으로 활용하고, Rogers(2003) 5속성을 T-P-S로 조작화하여 리터러시가 세 축을 연결하는 매개변수임을 이론화
국가 비교 연구	State Council of China (2017)	중국	정책 분석	정부주도 전략 프레임 제시	선행연구는 정책 문서. 본 연구는 이를 중국 모델 분류의 1차 자료로 활용하고, T-P-S 구조에서 P축(국가계획)이 T축(투자)·S축(순응)과 결합하는 정부주도형 확산 메커니즘으로 해석
	European Commission (2024)	EU	제도 분석	위험 기반 규제 체계 확립	선행연구는 규제 내용. 본 연구는 이를 규제중심형 모델의 근거로 활용하고, AI Act의 위험 기반 분류가 T축(기술 적용 속도)을 P축(제도)으로 관리하는 메커니즘임을 이론화
	김병운(2016), 정보화정책	한국·글로벌 분석	정책 동향 분석	국가AI전략·R&D 과제 제언	선행연구는 초기 정책 제언. 본 연구는 이를 한국 정책 프레임워크의 초기 조건 분석에 활용하고, 2019→2024 경로를 추적하여 혼합형 모델의 진화 과정을 실증

출처: 본 연구 재구성

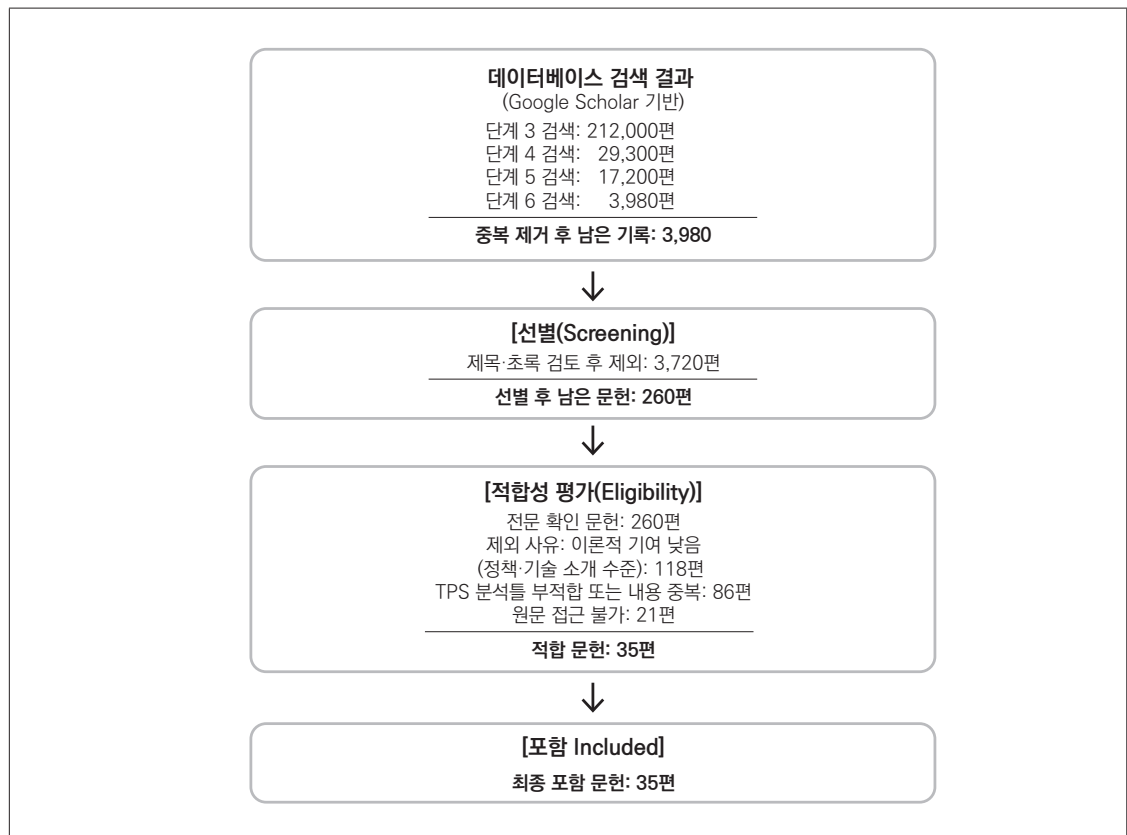
서비스 수용, 정책 의사결정, 국가전략 비교 등 개별 영역에서 중요한 통찰을 제공하였으나, 기술혁신-정책체계-사회적 수용을 하나의 분석틀로 통합하여 국가별 AI 확산 경로를 구조적으로 비교한 연구는 드물다. 이러한 연구 흐름을 체계적으로 정리하기 위해, 본 연구에서 검토한 주요 선행연구를 <표 2>와 같이 분류·요약하였다.

본 연구는 이러한 공백을 보완하고자 한다. T-P-S 삼각구조를 명시적 분석틀로 설정하고, 미국·EU·중국·한국의 확산 경로를 동일한 기준에서 비교함으로써 통합적 국가 비교 프레임워크를 제시한다는 점에서 기존 연구와 차별성을 갖는다.

III. 연구방법

본 연구는 2019년부터 2024년까지 발표된 정책 문서, 산업 보고서, 학술 논문을 대상으로 구조화된 내러티브 리뷰를 수행하였다. 이 시기는 생성형 AI 등장 이후 기술·정책·사회 전반에서 급격한 변화가 나타난 때로, 주요 국가들의 AI 전략이 본격적으로 구체화되기 시작한 전환점에 해당한다. 구조화된 내러티브 리뷰는 체계적 문헌고찰의 명시적 절차(검색·선정 기준)와 내러티브 리뷰의 맥락적 해석을 결합한 방법론이다. AI 확산을 기술혁신·정책체계·사회적 수용이라는 다층적 구조 속에서 이해하는 데 적합한 접근으로 판단하였다.

문헌 검색은 Web of Science(WoS), Scopus, Google



출처: 본 연구 작성 (PRISMA 2020 흐름도 기준)

<그림 1> PRISMA 2020 Flow Diagram

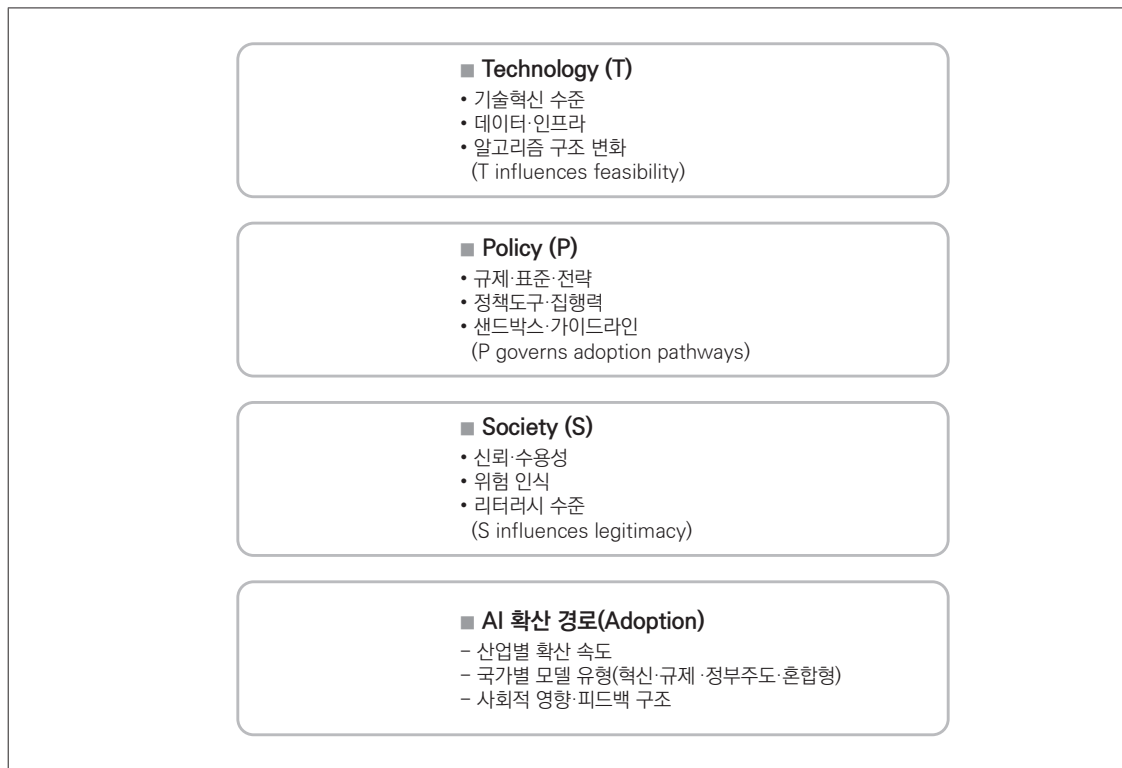
Scholar를 포함한 주요 학술 데이터베이스를 기반으로 수행하였다. WoS와 Scopus는 1차적으로 연구 범위와 핵심 주제어를 설정하는 데 사용하였으며, 실제 포괄적 수집과 정량적 PRISMA 절차는 검색 범위가 가장 넓은 Google Scholar를 중심으로 진행하였다.

영문 검색어는 ‘artificial intelligence’, ‘AI’, ‘machine learning’과 함께 ‘diffusion’, ‘adoption’, ‘implementation’, ‘policy’, ‘governance’, ‘regulation’ 등을 조합하여 구성하였다. 한국어 검색어는 ‘인공지능’, ‘정책’, ‘거버넌스’, ‘사회적 수용’, ‘디지털 리터러시’ 등을 사용하였다. Google Scholar 기반 검색에서는 검색식 정교화 단계를 거치며 총 212,000편에서 단계적으로 축소되었고, 중복 제거 후 3,980편이 남았다. 제목과 초록 검토 단계에서 3,720편을 제외하였으며, 전

문 검토 단계에서 분석틀 부적합과 기여도 부족, 원문 검증이 어려운 문헌 등의 사유로 225편을 추가 제외한 후 최종적으로 35편을 분석에 포함하였다. 전체 과정은 PRISMA 2020 흐름도에 따라 문서화하였다.

분석 단계에서는 기술혁신(Technology), 정책체계(Policy), 사회적 수용(Society)이라는 T-P-S 삼각구조를 중심으로 코딩 체계를 설계하였다. 오픈 코딩을 통해 핵심 개념을 도출한 뒤, 축 코딩을 통해 국가별·산업별·시기별 패턴을 연결하고, 선택 코딩을 통해 세 범주 간 상호작용을 설명하는 통합 해석틀을 구성하였다.

본 분석틀은 Rogers(2003)의 혁신확산이론과 Geels(2002)의 사회기술체계 이론을 접목하여 기술·정책·사회 요인을 구조적으로 연계하는 방식으로 재구성하였다. 혁신확산이론은 기술 수용의 단계적 과정과 영



출처: 본 연구 작성

〈그림 2〉 T-P-S 통합 개념도

향 요인을 설명하는 데 유용하며, 사회기술체계 이론은 제도와 맥락이 기술 변화에 미치는 영향을 포착하는 데 강점이 있다. 두 이론을 결합함으로써 국가별 정책 접근과 사회적 인식 구조가 기술 확산 경로에 미치는 영향을 체계적으로 해석할 수 있었다.

본 개념도는 기술혁신, 정책체계, 사회적 수용이 상호작용하며 국가별 AI 확산 경로를 형성하는 구조를 시각적으로 나타낸다.

타당성 확보를 위해 문헌 선정 기준, 제외 사유, 코딩 근거를 모두 기록하여 재현 가능성을 높이고자 하였다. 포함 여부가 모호한 문헌은 연구의 분석틀 적합성, 이론적 기여, 국가 비교 가능성을 기준으로 반복 검토하였다.

IV. 연구결과

1. 글로벌 AI 기술 전환과 산업별 확산

1) 기술 전환의 가속화

여러 보고서에서 AI 도입 속도의 가속화를 공통적으로 지적하고 있다. McKinsey(2023) 보고서는 생성형 AI가 연간 2.6조~4.4조 달러의 경제적 가치를 창출할 잠재력이 있으며, 특히 고객 운영, 마케팅 및 영업, 소프트웨어 엔지니어링, R&D 분야에서 파일럿 프로젝트를 활발하다고 분석한다. Stanford University(2024)의 AI Index도 유사한 경향을 보고하며, 2023년 한 해 동안 주요 AI 기업들의 투자 규모가 전년 대비 약 45% 증가했다고 밝혔다.

기술 전환은 산업별 데이터 인프라 구축 수준, 알고리즘 규모 확장, 모델 구조 혁신과 같은 기술적 요인(T)에 의해 주도되며, 특히 생성형 AI는 기존 기계학습 대비 학습 효율성과 범용성을 크게 확장시킴으로써 확산의 속도 자체를 변화시키는 기술적 분기점으로 작동한다.

다만 기술 전환이 확산에 미치는 영향은 산업마다 상당한 편차를 보인다. Bughin et al.(2018)의 분석에 따르면, AI 도입 선도 기업과 후발 기업 간 생산성 격차가 지속적으로 확대되고 있으며, 이는 '승자독식' 구조를

심화시킬 우려가 있다. 특히 중소기업의 경우 초기 투자 비용, 전문 인력 부족, 데이터 접근성 제약 등으로 인해 AI 도입이 지연되는 경향이 뚜렷하다.

동일한 기술혁신(T)이 실제 산업에서 적용되는 속도와 범위는 각국의 정책·규제 환경(P)에 따라 구조적으로 달라진다. 미국은 민간 자율성을 중심으로 신속한 실험을 허용하는 반면, EU는 AI Act를 통해 고위험 분야의 사전 규제와 감독을 강화함으로써 도입 속도를 제도적으로 관리하고 있다.

한국의 상황을 보면, 대기업과 중소기업 간 AI 활용 격차가 심각한 수준으로 나타나고 있다. 기술적·정책적 기반이 갖춰지더라도, 조직과 개인의 사회적 수용(S)과 디지털 리터러시 수준이 충분하지 않을 경우 확산은 지연된다. 특히 중소기업은 기술 불확실성과 책임 문제, 내부 전문성 부족(S)으로 인해 도입을 주저하는 경향이 강하다.

2) 산업별 확산 패턴

산업별 AI 확산 패턴을 질적으로 분석한 결과, 크게 선도(leading), 추격(catching-up), 후행(lagging)의 세 그룹으로 구분할 수 있었다.

선도 산업군(금융·IT)은 이미 갖춰진 디지털 인프라와 풍부한 데이터를 활용해 빠른 도입을 보이고 있다. 금융 분야에서는 문서 처리, 규제 대응, 리스크 관리 등에서 AI가 조직 프로세스에 깊이 통합되는 움직임이 관찰되며, 규제(P)가 강하게 작동하는 영역이라는 특성상 설명가능성·감사가능성 확보가 중요한 과제로 언급된다.

추격 산업군(제조·유통)은 초기 관망에서 적극 도입으로 전환하는 중이다. 제조업은 예지 정비와 품질 관리에서, 유통업은 수요 예측과 개인화 추천에서 성과가 나타나기 시작했다. 정부의 스마트팩토리 지원, 데이터 표준화, 세제 혜택 등 정책(P)은 기술혁신(T)이 실제 도입으로 이어지는 경로를 강화하는 역할을 한다.

후행 산업군(의료·교육)은 AI 도입 속도가 상대적으로 느린 편이다. 의료·교육 분야는 기술적 가능성(T)보다 사회적 수용(S)과 규제(P)의 영향력이 압도적으로

〈표 3〉 산업별 AI 확산 패턴(질적 개념들)

산업군	대표적용	확산단계	촉진요인	제약요인	대표사례
금융·IT (선도)	문서·코드생성, 위협/운영자동화	성숙	데이터풍부, ROI 가시성	규제준수복잡성	JP모건COiN, GitHub Copilot
제조·유통 (추격)	예지정비, 개인화추천	성장	운영효율, 고객경험	레거시, 데이터품질	GE Predix, 아마존추천
의료·교육 (후행)	진단보조, 개인화학습	도입	정확도향상가능성	안전, 윤리, 책임성	IBM Watson, Khan Academy

출처: 본 연구 재구성. 확산 단계는 McKinsey, 2023; Stanford University, 2024 등의 산업별 채택률 분석을 참고하여 질적으로 구분한 개념들이며, 정량적 입계치를 의미하지 않음.

* 확산 단계 정의는 Dwivedi et al., 2021; Bughin et al., 2018의 성숙도 연구를 토대로 질적 판단.

큰 영역이다. 의료에서는 진단 오류 가능성과 책임 소재 불명확성, 교육에서는 데이터 보호·공정성 우려가 확산을 구조적으로 제한한다.

은종환과 황성수(2020)는 AI 기반 정책 의사결정의 성패가 기술 수준보다 문제 구조화의 적절성에 좌우된다는 점을 사례 분석을 통해 제시하였다. 이는 정책체계(P)가 기술혁신(T)을 실제 현장에서 작동시키는 핵심 조건이라는 점을 뒷받침한다.

2. 국가별 정책 프레임워크와 확산 전략

본 절은 기술혁신(T), 정책체계(P), 사회적 수용(S)의 세 축을 공통 기준으로 삼아 미국, EU, 중국, 한국의 확산 전략을 비교한다. 네 국가의 ‘혁신주도형-규제중심형-정부주도형-혼합형’ 구분은 기술 자립도(T), 규제·거버넌스 구조(P), 사회적 신뢰·리터러시 기반(S)의 차이에 근거한 것으로, 경험적 자료에 기초한 유형화이다.

1) 미국: 혁신 주도형 모델

미국의 AI 확산은 기술혁신(T) 역량을 최우선으로 두는 정책철학(P)과, 민간 중심의 높은 기술 수용성(S)이 결합된 혁신주도형 모델로 특징지어진다. Stanford University(2024)에 따르면, 2023년 미국 AI 기업의 투자 규모는 672억 달러로 전년 대비 45% 증가했으며,

이 중 생성형 AI 분야가 252억 달러(37.5%)를 차지했다. McKinsey(2023)는 생성형 AI의 경제적 가치 창출 잠재력 \$2.6~4.4조 중 약 60%가 북미 시장에서 실현될 것으로 전망한다. White House(2023)의 행정명령은 안전성, 보안, 프라이버시 보호를 강조하면서도 혁신을 저해하지 않는 ‘light-touch’ 규제를 지향하며, 고위험 AI에는 안전 테스트·정보공개를 의무화하고 일반 응용은 자율규제를 허용한다.

미국 모델의 핵심은 민간 주도 혁신을 촉진하는 생태계 조성에 있다. 연방 정부는 R&D 투자, 인재 양성, 국제 표준화 참여를 통해 간접 지원하며, 규제는 최소한으로 유지한다. 다만 자율 규제 중심의 접근은 AI 안전성과 윤리적 우려에 대한 대응이 느리다는 비판을 받는다(Brundage et al., 2020).

2) EU: 규제중심형 모델

EU의 확산전략은 기술혁신(T)보다 규제와 기본권 보호(P)를 중심축으로 두고, 사회적 신뢰와 안전(S)을 정책의 최우선 가치로 설정하는 규제중심형 모델이다. Stanford University(2024)에 따르면, 2023년 EU의 AI 투자 규모는 약 78억 달러로 미국의 12% 수준에 그쳤으나, 규제 프레임워크 구축에서는 글로벌 선도적 위치를 차지한다. European Commission(2024)의 AI Act는 위험 기반 분류를 도입하여 AI 시스템을 금지·고위험·제한적 위험·최소 위험으로 구분하고, 각

단계에 차등적 규제를 적용함으로써 기술 확산의 방향과 속도를 제도적으로 관리한다.

고위험 AI(예: 생체 인식, 중요 인프라, 교육, 고용, 법 집행)는 엄격한 사전 승인, 투명성 요구, 인간 감독 의무를 준수해야 한다. EU 모델은 'Brussels Effect'로 불리는 규제의 글로벌 파급력을 활용하여, 사실상의 국제 표준을 선점하려는 전략을 구사한다. 다만 과도한 규제가 혁신을 저해하고, EU 기업들의 글로벌 경쟁력을 약화시킬 수 있다는 우려도 제기된다(Cath, 2018).

3) 중국: 정부 주도형 모델

중국은 국가전략 차원의 집중 투자(T), 강력한 정책 집행력(P), 광범위한 사회적 수용 기반(S)을 결합하여 정부주도형 확산모델을 구축하고 있다. Stanford University(2024)는 2023년 중국의 AI 투자가 약 134억 달러에 달했으나, 미국의 반도체 수출 통제로 인해 첨단 GPU 접근이 제약받고 있다고 분석한다. 최근 보고서들은 미국의 첨단 반도체 수출 통제가 중국의 AI 산업 발전에 중요한 변수로 작용하고 있으며, 이에 대응하여 중국이 자체 반도체 공급망과 기술 생태계 구축을 추진하고 있다고 분석한다. 「차세대 인공지능 발전 계획」(2017)은 2030년 세계 AI 혁신 중심지 도약을 목표로 삼고, 2020년 핵심 기술 확보, 2025년 주요 응용분야 선도, 2030년 종합적 세계 최고 수준 달성을 체계적으로 추진하고 있다.

중국 모델의 특징은 대규모 정부 투자, 산학연 통합, 데이터 집적, 그리고 강력한 정책 실행력이다. 다만 콘

텐츠 관리와 데이터 통제 중심의 정책 환경은 AI 산업의 개방성과 혁신 역량 측면에서 한계 요인으로 지적되며, 미국의 첨단 반도체 수출 통제로 인한 기술 접근성 제한 역시 중국 AI 산업의 구조적 리스크로 평가된다.

4) 한국: 혼합형 모델

한국의 확산전략은 기술혁신 촉진(T), 책임 기반 규제(P), 높은 사회적 기술수용성(S)을 조합한 혼합형 모델로 설명될 수 있다. 관계부처 합동(2019)의 「AI 국가전략」은 (1) AI 기술 경쟁력 확보, (2) 전 국민 활용 기반 조성, (3) 사람 중심의 AI 실현을 3대 축으로 제시하며 미국(혁신), EU(안전), 중국(정부주도)의 요소를 선택적으로 결합하고 있다. 과학기술정보통신부(2021)의 「신뢰할 수 있는 인공지능 윤리 기준」은 인간 존엄성 원칙, 사회의 공공선 원칙, 기술의 합목적성 원칙을 명시하며, 자율 준수를 권장하되 필요시 법제화 가능성을 열어두었다. 이는 EU의 강제적 규제와 미국의 자율 규제 사이의 중간적 접근으로 평가된다.

한국 모델의 강점은 높은 디지털 인프라, 우수한 IT 인력, 빠른 정책 실행력이다. 한국지능정보사회진흥원(2024) 조사에 따르면, 한국인의 AI 서비스 이용률은 상대적으로 높은 수준으로 나타났으며, 특히 20~30대에서 이용률이 두드러지게 높다. 신승윤(2024)은 OECD 38개국 분석에서 한국이 “AI 선도국” 그룹에 포함되지만, 기술혁신역량 지표에서는 미국·중국·EU 주요국보다 낮은 중위권에 위치한다고 지적했다. 다만 원천 기술 취약성, 대기업 편중, 데이터 접근성 제약, 글로벌

〈표 4〉 국가별 AI 정책프레임워크 비교

구분	미국	EU	중국	한국
정책 철학	시장 자율·혁신	권리·안전 중심	국가 주도	균형적 혼합
주요 도구	행정명령·가이드	AI Act (위험 기반)	국가 계획·투자	전략·윤리·샌드박스
관찰 패턴	빠른 혁신	신뢰 기반 구축	대규모 전개	점진적 확산

출처: White House(2023), European Commission(2024), State Council of China(2017), 관계부처합동(2019), 과학기술정보통신부(2021) 정책문서비교분석.

별 표준 영향력 부족은 한국 혼합형 모델의 구조적 과제로 지속적으로 제기된다.

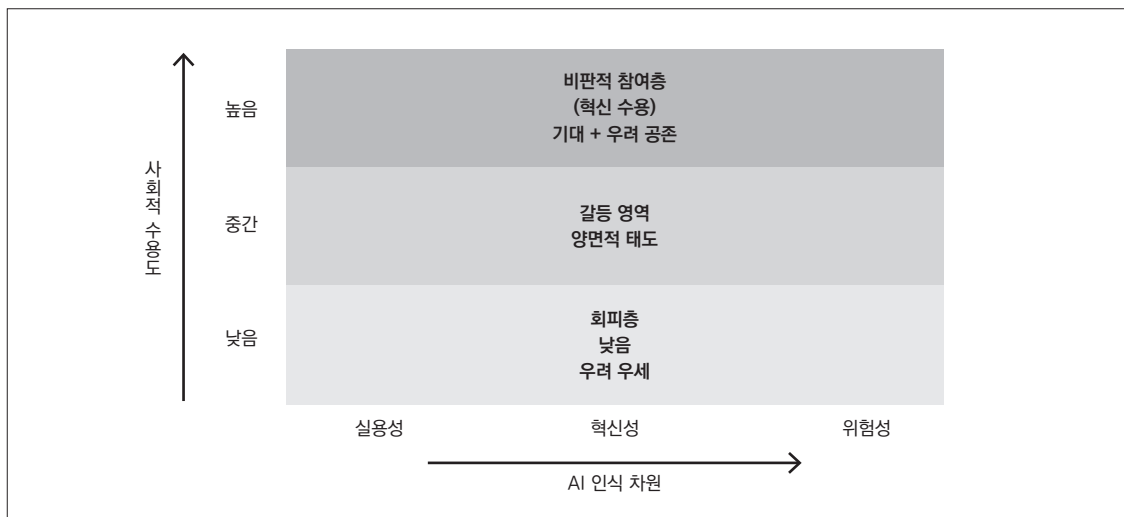
한국 내에서도 AI 확산은 균질하지 않으며, 조직 규모·역량·지역에 따른 구조적 불균등이 관찰된다. 대기업과 중소기업 간 AI 활용 격차는 기술 접근성뿐 아니라 전문 인력 확보, 데이터 인프라, 투자 여력의 차이에서 비롯되는 것으로 지적된다. 은종환과 황성수(2020)는 공공부문 AI 도입 사례 분석을 통해, 기술 자체의 성능보다 문제 구조화 수준과 조직의 정책 설계 역량이 성패를 결정한다는 점을 실증적으로 제시했다. 이는 혼합형 모델이 정책 문서 차원에서는 기술혁신(T)·정책체계(P)·사회적 수용(S)의 균형을 추구하나, 실제 확산 과정에서 다른 양상을 보임을 시사한다. 조직 단위의 역량 격차, 산업 부문 간 제도적 준비도 차이, 지역별 인프라 불균형이 확산 경로를 차별화하는 주요 요인으로 작용함을 시사한다. 따라서 한국의 혼합형 전략이 실효성을 갖기 위해서는 거시적 정책 방향 설정뿐 아니라, 중소기업·지역·공공부문의 AI 역량을 체계적으로 강화하는 미시적 지원 체계가 동시에 구축되어야 한다.

3. 사회적 수용 양상과 리터러시의 역할

1) 사회적 수용의 양면성

AI에 대한 사회적 수용은 기대와 우려가 공존하는 양면적 양상을 보인다. Edelman(2024)의 신뢰 바로미터 조사에 따르면, 글로벌 응답자의 52%가 AI 기술이 이익보다 해를 끼칠 것이라고 응답했으며, 특히 일자리 대체, 프라이버시 침해, 편향성 문제에 대한 우려가 높았다. 반면 동일한 조사에서 63%는 AI가 업무 효율성을 높이고 삶의 질을 개선할 것이라고 기대하며, 이러한 모순적 태도는 AI 확산의 사회적 맥락이 복잡함을 시사한다.

한국의 상황을 보면, 한국지능정보사회진흥원(2024)의 「2023 인터넷이용실태조사」에 따르면 국내 인터넷 이용자의 절반 이상이 AI 서비스를 이용한 경험이 있는 것으로 나타났으며, 특히 20~30대에서 상대적으로 높은 이용률을 보였다. 이는 AI 기술이 일상생활 전반으로 빠르게 확산되고 있음을 보여준다. 그러나 AI 기술의 활용 확대와는 별개로 개인정보 보호, 알고리즘 편향성, 책임성 등의 문제에 대한 사회적 우려 역시 지속적으로 제기되고 있다. 이러한 기대와 우려의 공존은



출처: 본 연구 재구성. 한국지능정보사회진흥원(2024) 및 Edelman(2024) 분석 결과를 바탕으로 구성. X축은 기술 수용도(낮음→높음), Y축은 우려 수준(낮음→높음)을 나타냄. 이는 정량적 임계치가 아니라 질적 개념틀임.

〈그림 3〉 한국 사회의 AI 수용 양면성 개념도

빠른 기술 확산 속에서 사회적 준비와 제도적 대응이 함께 요구되는 현실을 반영한다.

선행연구들은 AI에 대한 기대와 우려가 동시에 존재하는 양면적 수용 양상을 보고하고 있다. 일상생활의 편의성을 높이는 AI 서비스는 빠르게 수용되는 반면, 중요한 의사결정에 AI가 개입하는 것에 대해서는 신중한 태도를 보인다. 이는 성능 중심의 기술혁신(T)이 수용으로 자동 연결되지 않으며, 정책적 안전장치(P)와 사회적 신뢰(S)가 결합될 때 안정적 확산이 가능하다는 점을 시사한다.

2) AI 리터러시의 중요성

AI 리터러시는 AI 기술의 작동 원리를 이해하고, 비판적으로 평가하며, 윤리적으로 활용할 수 있는 역량을 의미한다. 단순히 AI 도구 사용법을 아는 것을 넘어, AI의 한계와 위험을 인식하고, AI가 생성한 정보의 신뢰성을 판단할 수 있는 메타인지 능력까지 포함한다. 장창기와 성옥준(2022)은 개인의 인식·신뢰·디지털 리터러시가 공공서비스 정책 수용 의도에 유의미한 영향을 미친다는 실증적 결과를 제시하며, 리터러시가 사회적 수용(S)을 구성하는 핵심 요인이라고 밝혔다.

University of Helsinki(2019)의 ‘Elements of AI’ 프로그램은 AI 리터러시 교육의 성공적 사례로 평가받는다. 이 무료 온라인 과정은 핀란드 인구의 약 1%인 55,000명 이상이 이수했으며, 현재는 40개 언어로 번역되어 전 세계 100만 명 이상이 수강했다. AI Singapore(2020)의 ‘AI for Everyone’ 프로그램도 유사한 접근을 취하며, 전 국민의 AI 역량 강화를 목표로 한다. 이러한 국가 차원의 리터러시 교육은 AI 확산의 사회적 기반을 다지는 핵심 전략으로 자리 잡고 있다.

AI 리터러시는 사회적 수용(S)의 핵심 구성요소이자, 기술혁신(T)과 정책체계(P)를 실제 확산으로 연결하는 매개 메커니즘으로 기능한다. 기술적 고도화가 빠르게 진행되더라도 시민이 AI의 작동 원리와 한계를 이해하지 못하면 신뢰 형성이 어렵고, 그 결과 알고리즘 책임성·데이터 이용 규범 등 정책체계(P)의 실효성 또한 확보되기 어렵다. 즉, AI 리터러시는 T-P-S 삼각구조를 연결하는 인지적 기반으로, 기술 발전과 정책 도입이 실제 현장에서 안정적으로 정착하는지를 결정하는 요인이다.

한국에서도 AI 리터러시 교육의 필요성이 강조되고 있으나, 아직 체계적인 국가 전략은 부재한 상황이다.

〈표 5〉 국가별 AI 확산 모델의 구조적 차이(T-P-S 기반 비교 요약)

구분	미국 (혁신주도형)	EU (규제중심형)	중국 (정부주도형)	한국 (혼합형)
T: 기술혁신 기반	빅테크 중심의 민간 R&D, 개방형 생태계, 고속 기술 선도	기술자립도 중간, 고위험 분야 중심 기술규제 병행	국가전략 투자·데이터 대규모 집적, 빠른 상용화	디지털 인프라 우수, 기술자립도 제한, 중소기업 격차 존재
P: 정책·규제 체계	최소규제(light-touch), 가이드라인 중심, 자율규제	AI Act 기반 강제 규제, 기본권·안전성 최우선	국가계획·정책 집행력 극대화, 통제 중심 규범	전략·윤리·샌드박스 병행, 미국·EU 요소 혼합
S: 사회적 수용 기반	시장 중심 혁신 수용성 높음, 위험 인식 낮음	안전·투명성 요구 높음, 규범 준수 기대 강함	국가 신뢰 기반 수용성 높음, 규제 순응도 높음	이용 경험 확대, 동시에 편향·책임성 우려 공존 (양면성)
확산 경로의 특징	민간혁신 → 시장확산 주도	규제설계 → 신뢰 기반 확산	국가전략 → 대규모·신속 확산	균형적 확산이나 기술·생태계 구조적 제약 존재
정책적 함의	혁신 촉진 중심 모델의 한계는 안전성·책임성의 외부화	규제 기반이 확산의 안정성을 확보하나 혁신 속도 제한	정부주도 확산은 속도·규모 장점, 기술·표준 의존도 위험	혼합형 모델의 전략적 정교화 필요-일관된 정책철학 요구

출처: 연구 재구성. White House(2023), European Commission(2024), State Council of China(2017), 관계부처 합동(2019) 등 비교 분석

한국지능정보사회진흥원(2024)은 AI 리터러시 교육이 기술 중심 교육에서 벗어나, 윤리적 판단력, 비판적 사고력, 사회적 책임감을 함께 기르는 방향으로 설계되어야 한다고 제안한다. AI 리터러시 정책은 기술 활용 교육을 넘어 AI의 한계와 위험, 윤리적 책임을 이해하는 시민 역량 강화 방향으로 발전할 필요가 있다.

V. 결론

1. 연구결과 요약

본 연구는 인공지능 확산을 기술혁신, 정책체계, 사회적 수용의 삼각 구조로 파악하고, 이 분석틀을 활용하여 미국, EU, 중국, 한국의 확산 경로를 비교 분석하였다.

첫째, AI 기술은 규칙 기반 전문가 시스템에서 기계 학습, 심층학습, 대규모 언어 모델에 이르는 패러다임 전환을 거쳤으며, 특히 2020년 이후 생성형 AI의 등장 이 확산의 분기점이 되었다.

둘째, 산업별 AI 확산은 디지털 인프라, 데이터 가용성, 규제 환경에 따라 선도(금융·IT), 추격(제조·유통), 후행(의료·교육) 산업군으로 구분된다.

셋째, 국가별 정책 프레임워크는 AI 확산 경로를 구조적으로 제약하는 핵심 매개 요인이다. 미국(혁신 주도형), EU(규제 중심형), 중국(정부 주도형), 한국(혼합형) 모델은 상이한 정책 철학을 바탕으로 기술혁신과 사회적 수용을 서로 다르게 연결한다. 한국의 혼합형 모델은 유연성을 가지지만 중장기 방향성 확립이라는 과제를 안고 있다.

넷째, AI에 대한 사회적 수용은 기대와 우려가 공존한다. 한국은 높은 디지털 인프라와 이용률에도 불구하고 편향, 프라이버시, 일자리 대체에 대한 우려가 상당하며, 이는 사회적 신뢰 기반이 필수적임을 시사한다.

다섯째, 핀란드와 싱가포르 사례는 국가 차원의 체계적 AI 교육이 사회적 수용 기반을 강화하고, 기술혁신과 정책체계를 현실에서 구현하는 핵심 인프라로 기능함을 보여준다.

2. 이론적 및 실증적 기여

이 연구의 이론적 기여는 다음 세 가지로 정리할 수 있다. 첫째, Rogers(2003)의 혁신확산이론과 Geels(2002)의 사회기술체계 이론을 AI 맥락에 맞게 통합하여 T-P-S 분석틀을 제시했다. 둘째, AI 확산을 정태적 현상이 아니라 기술·정책·사회가 함께 변화하는 공진화 과정으로 분석하여 동태적 특성을 포착했다. 셋째, 4개국 비교를 통해 한국 혼합형 모델의 강점과 한계를 구조적으로 제시하여 후발국의 전략 선택에 이론적 근거를 제공했다.

실증적으로는 정책 프레임워크를 기술과 사회를 연결하는 핵심 매개 구조로 재구성하여 정책 맥락의 중요성을 부각했고, 국내 정보화정책 연구를 통합적으로 활용해 한국의 혼합형 전략이 실제 확산 경로에서 어떻게 구현되는지를 사례 중심으로 보여주었다.

3. 정책적 시사점

T-P-S 삼각구조에 기반한 비교 분석 결과, 한국 AI 거버넌스 발전을 위해 다음과 같은 다섯 가지 정책적 시사점을 도출하였다.

첫째, 위험 기반 규제 체계의 정교화가 필요하다. EU AI Act의 위험 기반 분류를 참고하되, 규제 샌드박스를 정책 학습의 실험 공간으로 재정립하고 위험 평가 기준을 투명하게 공개해야 한다.

둘째, 속의 거버넌스의 실질화가 필요하다. 고위험 AI 도입이나 공공부문 알고리즘 활용 시 시민·산업계·전문가가 참여하는 상설 대화 구조와 속의형 여론 수렴 절차를 제도화해야 한다.

셋째, 포용적 AI 리터러시 정책을 확립해야 한다. 단순 기술 교육을 넘어 알고리즘의 한계와 편향, 윤리적 쟁점을 이해하는 시민역량 강화 전략으로 설계하고, 대상별 맞춤형 프로그램을 마련해야 한다.

넷째, 중소기업과 공공부문의 AI 활용 역량을 체계적으로 강화해야 한다. 단순 보급을 넘어 문제 정의-데이

터 정비-조직 프로세스 재설계를 포괄하는 통합 지원 체계를 장기 전략으로 추진할 필요가 있다.

다섯째, 국제 규범 논의에서의 능동적 역할을 강화해야 한다. OECD, ISO, IEEE 등 국제기구 논의를 체계적으로 모니터링하고, 공공서비스와 리터러시 정책에서 축적된 경험을 제안하여 사람 중심·신뢰 기반 AI라는 한국형 가치를 확산시켜야 한다.

4. 연구의 한계와 향후 과제

본 연구는 기술혁신-정책체계-사회적 수용의 삼각구조를 기반으로 주요 국가의 AI 확산 경로를 비교·분석하였다. 다만 다음과 같은 한계가 존재한다.

첫째, 구조화된 내러티브 리뷰는 정책 맥락 이해에 유용하지만 체계적 문헌고찰(PRISMA)의 엄격성이나 정량적 메타분석을 완전히 대체하지 못한다. 문헌 선정 과정에서 연구자 판단이 개입될 여지가 있고, 정책 문서와 학술 연구 간 개념 차이를 재구조화하는 과정에서 일부 정보 단순화가 불가피했다.

둘째, 분석 시계열이 2019~2024년으로 제한되어 2023년 이후 급변하는 생성형 AI 정책 환경을 충분히 반영하지 못했다. 또한 영어·한국어 문헌 중심 분석으로 일본·독일·프랑스 등 중견국의 정책 철학이 충분히 포함되지 못했다.

셋째, 기술-정책-사회의 삼각구조는 국가 간 비교에 유용하지만 각국 제도·문화·산업 구조의 복잡성을 모두 포착하기 어렵다.

넷째, 미·중 기술패권 경쟁, 글로벌 공급망 재편, 플랫폼 기업 전략 변화 등 외생적 정치경제 요인을 단일 분석틀로 충분히 설명하기 어려웠다.

향후 연구에서는 AI 투자 규모, 규제 강도, 사회적 신뢰 수준 등을 포함한 정량·정성 혼합연구와 구조방정식 모형(SEM), 정책네트워크 분석 등 복합적 방법론 적용이 요구된다. 또한 공공부문 AI 활용 사례, 산업 유형·조직 규모에 따른 도입 격차, 국제 규범·표준화 논의의 장기적 영향을 심층적으로 다루는 후속연구가 요구된다.

■ 참고문헌

- 관계부처합동 (2019). <인공지능 국가전략>. 서울: 관계부처합동.
- 과학기술정보통신부 (2021). <신뢰할 수 있는 인공지능 개발 안내서>. 과천: 과학기술정보통신부.
- 김병운 (2016). 인공지능 동향분석과 국가차원 정책제언. <정보화정책>, 23권 1호, 74-93.
- 신승운 (2024). 국가 AI 경쟁력에 따른 OECD 국가 유형 분류: 퍼지셋 이상형 분석을 중심으로. <정보화정책>, 31권 2호, 39-64. <https://doi.org/10.22693/NIAIP.2024.31.2.039>
- 장창기·성옥준 (2022). 인공지능 기반 공공서비스 정책수용 의도에 관한 연구: 개인의 인식과 디지털 리터러시 수준이 미치는 영향을 중심으로. <정보화정책>, 29권 1호, 60-83. <https://doi.org/10.22693/NIAIP.2022.29.1.060>
- 은종환·황성수 (2020). 인공지능을 활용한 정책의사결정에 관한 탐색적 연구: 문제구조화 유형으로 살펴 본 성공과 실패 사례 분석. <정보화정책>, 27권 4호, 47-66. <https://doi.org/10.22693/NIAIP.2020.27.4.047>
- 한국지능정보사회진흥원 (2024). <2023 인터넷이용실태조사>. 대구: 한국지능정보사회진흥원.
- 한국지능정보사회진흥원 (2024). <2024 지능정보윤리 이슈리포트 겨울호(Vol. 5, No. 4)>. 대구: 한국지능정보사회진흥원.
- AI Singapore (2020). *AI for Everyone Programme*. Singapore: National AI Programme Office.
- Brown, T., Mann, B., Ryder, N., Subbiah, M., Kaplan, J., Dhariwal, P.,
- Neelakantan, A., Shyam, P., Sastry, G., Askell, A., Agarwal, S.,
- Herbert-Voss, A., Krueger, G., Henighan, T., Child, R., Ramesh, A.,
- Ziegler, D. M., Wu, J., Winter, C., ... Amodei, D. (2020). Language models are few-shot learners. *Advances in Neural Information Processing Systems*, 33, 1877-1901.
- Brundage, M., Avin, S., Wang, J., Belfield, H., Krueger, G., Hadfield, G., Khlaaf, H., Yang, J., Toner, H., Fong, R., Maharaj, T., Koh, P. W., Hooker, S., Leung, J., Trask, A., Blumke, E., Lebensold, J., O'Keefe, C., Koren, M., Ryffel, T., Rubinovitz, J.

- B., Besiroglu, T., Carugati, F., Clark, J., Eckersley, P., de Haas, S., Johnson, M., Laurie, B., Ingerman, A., Krawczuk, I., Askell, A., Cammarota, R., Lohn, A., Krueger, D., Stix, C., Henderson, P., Graham, L., Prunkl, C., Martin, B., Seger, E., Zilberman, N., Ó hÉigeartaigh, S., Kroeger, F., Sastry, G., Kagan, R., Weller, A., Tse, B., Barnes, E., Dafoe, A., Scharre, P., Herbert-Voss, A., Rasser, M., Sodhani, S., Flynn, C., Gilbert, T. K., Dyer, L., Khan, S., Bengio, Y. & Anderson, H. (2020). *Toward trustworthy AI development: Mechanisms for supporting verifiable claims*. arXiv:2004.07213.
- Buchanan, B. G. & Shortliffe, E. H. (1984). *Rule-based expert systems: The MYCIN experiments of the Stanford Heuristic Programming Project*. Reading, MA: Addison-Wesley.
- Bughin, J., Seong, J., Manyika, J., Chui, M. & Joshi, R. (2018). *Notes from the AI frontier: Modeling the impact of AI on the world economy*. McKinsey Global Institute.
- Cath, C. (2018). Governing artificial intelligence: Ethical, legal and technical opportunities and challenges. *Philosophical Transactions of the Royal Society A*, 376(2133), 20180080. <https://doi.org/10.1098/rsta.2018.0080>
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319-340. <https://doi.org/10.2307/249008>
- Dwivedi, Y. K., Hughes, L., Ismagilova, E., Aarts, G., Coombs, C., Crick, T., Duan, Y., etc. (2021). Artificial Intelligence (AI): Multidisciplinary perspectives on emerging challenges. *International Journal of Information Management*, 57, 101994. <https://doi.org/10.1016/j.ijinfomgt.2019.08.002>
- Edelman (2024). *Edelman Trust Barometer 2024: Technology and Trust*. New York: Edelman.
- European Commission (2024). Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act).
- Geels, F. W. (2002). Technological transitions as evolutionary reconfiguration processes. *Research Policy*, 31(8-9), 1257-1274. [https://doi.org/10.1016/S0048-7333\(02\)00062-8](https://doi.org/10.1016/S0048-7333(02)00062-8)
- Krizhevsky, A., Sutskever, I. & Hinton, G. E. (2012). ImageNet classification with deep convolutional neural networks. *Advances in Neural Information Processing Systems*, 25, 1097-1105.
- McKinsey (2023). *The economic potential of generative AI*. New York: McKinsey & Company.
- OECD (2023). The State of Implementation of the OECD AI Principles Four Years On. *OECD Digital Economy Papers*, 359. <https://doi.org/10.1787/835641c9-en>
- Page, M. J., McKenzie, J. E., Bossuyt, P. M., Boutron, I., Hoffmann, T. C., et al. (2021). The PRISMA 2020 statement. *BMJ*, 372, n71. <https://doi.org/10.1136/bmj.n71>
- Radford, A., Wu, J., Child, R., Luan, D., Amodei, D. & Sutskever, I. (2019). Language models are unsupervised multitask learners. OpenAI Blog. <https://openai.com/index/better-language-models>
- Rogers, E. M. (2003). *Diffusion of innovations (5th ed.)*. New York: Free Press.
- Stanford University (2024). *Artificial Intelligence Index Report 2024*. Stanford, CA.
- State Council of the People's Republic of China (2017). *New Generation Artificial Intelligence Development Plan*. Beijing: State Council.
- University of Helsinki (2019). *Elements of AI: A free online course*. Helsinki: University of Helsinki & Reaktor.
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł. & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30, 5998-6008.
- White House (2023). *Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence*. Washington, D.C.