

석사학위논문

기술수용모델을 활용한 AI 기술이
사용의도에 미치는 영향에 관한 연구
- AI 리터러시의 매개효과 검증을 중심으로 -

2025년

한성대학교 지식서비스&컨설팅대학원

지식서비스&컨설팅학과

ESG 경영 컨설팅 전공

윤 종 현

석사학위논문
지도교수 정진택

기술수용모델을 활용한 AI 기술이 사용의도에 미치는 영향에 관한 연구

- AI 리터러시의 매개효과 검증을 중심으로 -

A Study on the Impact of Artificial Intelligence Technology on
Intention to Use Using Technology Acceptance Model

- Focused on the Mediating Effects of AI Literacy -

2024년 12월 일

한성대학교 지식서비스&컨설팅대학원

지식서비스&컨설팅학과

ESG경영컨설팅전공

윤 종 현

석사학위논문
지도교수 정진택

기술수용모델을 활용한 AI 기술이 사용의도에 미치는 영향에 관한 연구

- AI 리터러시의 매개효과 검증을 중심으로 -

A Study on the Impact of Artificial Intelligence Technology on
Intention to Use Using Technology Acceptance Model

- Focused on the Mediating Effects of AI Literacy -

위 논문을 컨설팅학 석사학위 논문으로 제출함

2024년 12월 일

한성대학교 지식서비스&컨설팅대학원

지식서비스&컨설팅학과

ESG경영컨설팅전공

윤 종 현

윤종현의 컨설팅학 석사학위 논문을 인준함

2024년 12월 일

심사위원장 주 형 근 (인)

심 사 위 원 이 형 용 (인)

심 사 위 원 정 진 택 (인)

국문초록

기술수용모델을 활용한 AI 기술이
사용의도에 미치는 영향에 관한 연구
- AI 리터러시의 매개효과 검증을 중심으로 -

한성대학교 지식서비스&컨설팅대학원

지식서비스 & 컨설팅학과

E S G 경영 컨설팅 전공

윤 종 현

최근 생성형 AI인 ChatGPT의 출현 이후 전세계적으로 인공지능(Artificial Intelligence) 기술이 기업, 학교, 가정 등 다방면의 분야에서 확산하고 있으며, 급속도로 진화하는 AI 기술로 인해 긍정과 부정 영향이 동시에 나타나고 있다. 인간과 인공지능 기술이 공존하게 될 사회에서 AI 기술에 대한 이해와 활용은 사용자 측면에서 필수적인 요소가 될 것이다. AI가 제공하는 정보에 대해 주체적, 비판적으로 평가하면서 AI 기술과 소통할 수 있는 AI 리터러시에 대한 중요성이 높아지고 있다. 선행연구들은 AI가 가지고 있는 기술적 영향력에 대한 분야가 주로 이루어지고 있으며, AI 기술 사용에 있어 AI 리터러시와의 영향력 관계를 살펴본 연구는 제한적이다.

이에 본 연구에서는 새로운 기술과 서비스를 받아들이는 이용자들의 기술 수용 과정 및 사용의도를 파악하기 위해 기술수용모델을 기반으로

설명하고 있다. 또한 AI 기술에 대한 이해도, 활용 능력 및 비판 역량으로 대표되는 AI 리터러시 수준과 사용의도 간의 관계에 관해 설명하고자 하였다. 구체적으로 AI에 관한 기술수용모델의 변수로 지각된 유용성과 사용용이성을 설정하였고 AI 리터러시라는 요인을 매개변수로 설정하여 이용자가 AI 기술을 사용하려는 의도에 영향을 미칠 것인지 구조방정식 모델을 통해 검증하고자 하였다.

본 연구는 AI 기술에 대한 사용자들의 수용 인식을 파악하기 위하여 2024년 12월에 만 20세 이상 성인 312명을 대상으로 실시한 온라인 설문조사 결과를 기반으로 분석하였다. AI 기술에 대한 조사 데이터를 활용하여 실증분석을 수행하였다. 주요 연구 결과는 다음과 같다.

기술수용모델의 주요 변수 중 사용용이성은 지각된 유용성과 사용의도에 유의미한 영향을 미치는 것으로 나타났으며 AI 리터러시에는 가설과 같이 정(+)의 영향을 미치는 것으로 나타났다. 지각된 유용성은 AI 리터러시 수준과 사용의도에 정(+)의 영향을 미치는 것으로 나타났다. AI 리터러시 수준도 사용의도에 정(+)의 영향을 미치는 것으로 나타났다. 그리고 사용용이성과 지각된 유용성이 사용의도에 영향을 미치는 데 있어 AI 리터러시 수준이 매개 역할을 통한 간접효과를 일으키는 지 확인 결과 모두 유의미한 영향을 미치는 것으로 확인되었다.

본 연구는 AI 기술의 사용의도에 영향을 미치는 요인 중 그 중요성이 대두되고 있는 AI 리터러시를 기존의 이론적 모델의 적용 및 확장을 시도하고 실증적으로 검증하였다는 점에서 의의가 있으며, 사용자가 AI 기술을 수용하는 데 AI 리터러시 수준이 매개효과가 있는 것을 확인하였다는 점에서 의의가 있다. 또한 향후 급속한 AI 기술의 발전이 예상됨에 따라 예상되는 AI 기술 사용자의 양극화 이슈 및 AI 리터러시 격차 해소를 위한 정책 수립 시 본 연구의 실증 결과를 기초 자료로 제공할 수 있다는 점에서 학문적 의의가 있다.

이러한 연구결과를 바탕으로 첫째, AI 기술 발전이 지속됨에 따라 사회 계층별 새로운 형태의 AI 리터러시 격차는 지속적으로 커질 것으로 예상된다. 이를 해소하기 위해 단순한 정보기기 활용 역량에 대한 리터러시

가 아닌 더욱 포괄적 측면에서 비판적 AI 리터러시 역량 강화가 필요하다. AI 리터러시 역량 강화와 함께 AI 리터러시 격차 해소를 위해 다양성, 형평성과 포용성(DE&I: Diversity, Equity, & Inclusion) 개념이 필수적으로 전제되어야 한다.

둘째, 현재 일상생활에서 남녀노소 불편함 없이 사용 중인 휴대전화처럼 AI 기술 사용을 확대하기 위해 AI의 사용용의성과 지각된 유용성에 관한 긍정적 인식 확산이 필요하다. 특히, 전체 국민 3명 중 2명은 고용 및 노동 불안정 문제를 AI 기술로 인해 가장 악화할 것으로 예상하며 걱정하고 있다. 부정적 인식의 극복을 위해 사회안전망의 합의와 구축, 노동자의 재교육 시스템 등 국가, 기업, 국민의 사회적 합의를 통한 적극적인 대응책이 요구된다.

셋째, 개인 정보 보호와 AI 기술의 오류에 대한 문제를 해결하여 AI 사용자의 신뢰도와 활용도를 개선할 방안을 모색해야 한다. AI 알고리즘의 투명성을 높여 사용자들이 AI의 작동 방식을 이해할 수 있도록 하고 AI의 학습 데이터에 대한 검증과 평가를 강화하여 AI 기술의 성능과 정확도를 높여야 한다. 또한, AI 기술의 부작용을 예방하고 대처할 수 있는 법적 제도와 규제를 마련하여 사용자들의 권리와 안전을 보장할 필요가 있다.

본 연구는 설문조사 데이터가 한국 내 사용자에게 한정되어 있다는 점, 설문이 특정 시점에 이루어져 시간에 따른 인식과 행동 패턴 변화를 완전히 반영하지 못할 수 있다는 점, 설문 문항과 데이터를 온라인 설문조사를 통해 제한된 계층 인원들을 표본으로 활용함에 따른 AI 기술 관련 모든 영향 변수를 반영하지 못할 수 있다는 연구의 한계가 있다. 향후 AI 기술 수용성에 관련된 연구 방향으로 법적·윤리적 사회문제를 적용한 AI 리터러시 개념, 역량 및 활용에 대한 연구 등이 수행될 수 있을 것으로 기대한다.

【주요어】 인공지능(AI), AI 리터러시, 기술수용모델(TAM), 지각된 유용성, 사용용의성, 사용의도

목 차

제 1 장 서론	1
제 1 절 연구배경 및 목적	1
제 2 절 연구방법 및 범위	4
제 3 절 연구의 구성	6
제 2 장 이론적 배경	8
제 1 절 AI 기술	8
제 2 절 기술수용모델	12
제 3 절 사용의도	14
제 4 절 지각된 유용성	15
제 5 절 사용용이성	16
제 6 절 AI 리터러시	18
제 3 장 연구설계	29
제 1 절 연구모형 및 연구가설 설정	29
제 2 절 변수의 조작적 정의	32
제 3 절 설문자료 수집	33
제 4 절 설문지 구성	34

제 5 절 분석방법	36
 제 4 장 연구결과	 37
제 1 절 인구통계학적 특성 및 기술통계 분석	37
제 2 절 신뢰도 및 타당도 분석	41
제 3 절 주요변수의 상관관계 분석	43
제 4 절 확인적 요인분석	44
제 5 절 연구가설 검증	49
 제 5 장 결론	 58
제 1 절 연구결과	58
제 2 절 시사점	59
제 3 절 한계 및 향후 연구	62
 참 고 문 헌	 64
 부 록	 77
 ABSTRACT	 83

표 목 차

[표 2-1] 리터러시 개념 변화	19
[표 3-1] 변수의 조작적 정의 및 선행연구	32
[표 3-2] 설문 조사 개요	33
[표 3-3] 통계 분석 방법	36
[표 4-1] 표본의 인구통계학적 특성	38
[표 4-2] 주요 변수의 측정 항목별 평균 및 표준편차	39
[표 4-3] 측정변수의 기술통계 분석	40
[표 4-4] 신뢰도 및 타당도 분석결과	42
[표 4-5] 상관관계 분석	43
[표 4-6] 확인적 요인분석 모형적합도 분석	45
[표 4-7] 확인적 요인분석 결과: 집중타당도	47
[표 4-8] 구성개념신뢰도(CR)와 평균분산추출(AVE)	48
[표 4-9] 연구모형 모형적합도 분석	49
[표 4-10] 구조모형 분석 결과	53
[표 4-11] 경로 간접효과	55
[표 4-12] 표준화 직접효과, 간접효과 및 총효과	56
[표 4-13] 연구가설 검정 결과 요약	57

그 림 목 차

[그림 3-1] 연구모형	29
[그림 4-1] 연구모형 검정 결과	52

제 1 장 서론

제 1 절 연구배경 및 목적

2022년 11월 OpenAI가 ChatGPT-3를 발표했다. 2023년 11월에는 이전 모델 보다 최신의 정보(2023년 4월까지)를 처리할 수 있는 GPT-Turbo가 발표되어 훨씬 정교하고 다양한 작업이 가능해졌다(이은창, 2023). 그리고 2024년 1월 10일 OpenAI는 GPT 스토어를 공개했다. GPT 스토어는 기업·개인이 OpenAI의 GPT 모델 기반으로 개발한 맞춤형 챗봇을 거래할 수 있는 플랫폼이다(파이낸셜뉴스, 2024). 구글, 애플의 애플리케이션 마켓처럼, OpenAI가 인공지능(Artificial Intelligence, 이하 AI) 생태계를 장악할 것으로 예상된다(헤럴드경제, 2024). 2024년 5월에는 OpenAI가 GPT-4o 신모델을 출시하였으며, ChatGPT가 메모리 기능을 갖추어 사용자와의 이전 대화를 통해 학습하고 실시간 번역을 하는 등 사용 편의성 측면에서 진전을 이루었다(이데일리, 2024). 2024년 9월에는 추론에 특화된 새로운 AI 모델인 OpenAI o1(오원)을 출시했다. 이 모델은 반응하기 전에 생각하는 데 더 많은 시간을 할애하도록 설계, 복잡한 작업을 추론하고 과학, 코딩, 수학 분야의 이전 모델보다 더 어려운 문제를 해결할 수 있는 것으로 알려졌다(한국일보, 2024). 신모델 출시 속도가 더욱 빨라지고 있으며 전문가가 아닌 일반인들은 이제 AI 기술의 진화 속도를 따라가기 버거운 상황이다.

ChatGPT 이용자 수는 23년 12월 기준 한 달간 16억 명가량으로 나타났다. 2023년 1월 1억 명, 5월에는 18억 명으로 정점을 찍은 바 있다(Similarweb, 2024). 비록 ChatGPT에 관한 관심이 전세계적으로 달아올랐던 열기가 식었고 일시적인 화제성도 드물어지는 추세이지만 이제는 AI 기술이 일상 곳곳으로 영향을 미치고 있다. 스마트폰, 노트북, PC 및 가전 제품에도 필수 기능으로 빠질 수 없는 요소가 되고 있다. 2024년 1월 9일 개최된 CES에 관통하는 주제는 AI 기술이었다. 반도체에서 로봇에 이르

기까지, AI 생태계는 커지고 있다. 마이크로소프트는 “2024년은 AI PC의 해가 될 것”이라고 밝힌 바 있다(조선일보, 2024).

매일 논문을 공부하는 연구자에게 ChatGPT는 시간 절약 도구이다. 수백 페이지 자료를 한 번에 한 장짜리 요약으로 바꿔 준다. 대학가에서는 필수 학습 도우미이고 수강하려는 수업의 개요를 소개해 주거나 프로그래밍 과제를 위한 코딩도 대신 써 준다. 성균관대 교육개발센터가 학생 415명을 대상으로 2023년 11월 실시한 설문조사에서 ChatGPT를 비롯한 생성형 AI의 조언에 만족한다는 답변 비율은 81.9%에 달했다(이상은, 김예진, 구민영, 이수민, 2023). 반면 AI가 학습 보조 도구로서 쓸모가 없다는 반응은 열 명 중 한 명꼴에 불과했다. 마이크로소프트(MS)는 워드, 엑셀, 파워포인트로 구성된 사무용 소프트웨어에 ChatGPT 기반의 생성형 AI 기능을 넣었다. 보고서와 발표자료 초안을 기계가 수월하게 만든다. 할리우드에서는 생성형 AI가 영화, 드라마의 스토리 기획안과 대본 초고를 여러 번으로 무한정 생산해 내고 있다. AI 작가의 비중이 급증해 인간 일자리카지 대체될 조짐이 보이자, 미국의 영상 작가 조합은 2023년 140여 일의 이례적 파업을 벌였다(Associated Press, 2023).

미국 증시를 주도하는 ‘매그니피센트7’ 즉, 애플, 마이크로소프트, 알파벳(구글), 아마존, 메타(페이스북), 테슬라, 엔비디아 등 7대 기술 우량주의 성장 동력은 모두 생성형 AI이다(주간동아, 2023). 생성형 AI가 휴대폰, 전기나 인터넷과 같은 보편 기술이 될 것이라는 전망은 빠르게 현실이 되고 있다.

2024년 1월 15일 개막된 다보스포럼에서 글로벌 105개국 CEO 4,702명 대상 설문조사를 진행한 결과, 58%는 1년 내 AI가 제품의 서비스 품질을 향상할 수 있다고 응답했으며, 70%는 향후 3년 내 AI가 기업의 일하는 방식을 근본적으로 바꿀 것이라고 내다봤다. 25%는 AI 도입에 따라 올해 인력을 5% 감축할 계획이라고 답했다(PwC, 2024). 2024년 1월 14일 국제통화기금(IMF)은 ‘AI와 일의 미래’라는 보고서를 발표하였다. AI 기술이 발전할수록 전 세계의 일자리 중 약 40%가 영향을 받을 가능성이 있으며, 사회 불평등이 심화할 것이라는 연구 결과를 제시하였다. “역사적

으로 자동화와 정보기술의 발전은 일상적이고 반복적인 일에 영향을 미쳤지만, AI는 고학력·고숙련 노동자의 일자리에 충격을 준다는 점에서 구별된다”라고 밝혔다(IMF, 2024). 이런 특성으로 인해 AI로 인한 일자리 감소는 선진국이 개발도상국보다 영향이 더 클 것으로 전망됐다. 선진국의 경우 일자리의 60% 정도가 영향을 받을 것으로 조사됐다. 그중 절반은 AI로 인해 업무 통합과 생산성 향상으로 혜택을 누릴 것으로 보이지만, 나머지 절반은 여전히 인간이 하는 일의 주요 업무가 AI로 대체될 것으로 예상된다. 후자의 경우 노동 수요가 축소되어 임금이 감소하고, 심할 경우 일자리 자체가 통째로 사라질 가능성이 크다. 신흥 성장 국가에서는 일자리의 40%가 AI의 영향을 받게 되고, 저소득 국가에서는 26%로 영향력이 낮아질 것으로 조사되었다. 저소득 국가들은 경제발전 수준이 낮아 AI의 직접 영향을 적게 받을 것으로 보인다. 그러나 사실은 이러한 국가들이 AI를 이용한 경제적 혜택을 이끌어낼 만한 숙련된 인력과 인프라가 준비되어 있지 않다는 의미이다. 따라서 시간이 지날수록 AI를 이용해 앞서나가는 선진국과의 격차가 더 커지고, 이는 국가 간 불평등을 더 악화하는 요인이 될 수 있다고 보고서에서 강조되었다(IMF, 2024).

AI 기술의 급격한 발전은 인간사회에 빛과 그늘을 양산하게 될 것이다. AI 기술은 사회의 많은 문제들을 개선하는 동인이 되어 긍정적인 영향을 미친다. 한편, 이미 부정적인 사회문제들도 동시에 발생하고 있다. 온라인상에서 AI로 인해 생성된 딥페이크(Deepfake) 이미지와 가짜뉴스 등 거짓 정보가 확산되는 등 위험 요소가 상존하고 있다. AI가 딥러닝(Deep Learning)과 같은 학습기능을 통해 정교하게 진화를 거듭하게 되면서 진짜와 가짜를 구분하기가 더욱 어려워지고 있다. ‘블랙박스’로 비유되는 알고리즘도 여러 부작용을 양산하고 있고, 콘텐츠 추천 기능이 광고에서 광범위하게 확산하고 있다. 이에 따라 소비자는 자신도 모르는 사이에 알고리즘의 편향성에 빈번하게 갇히게 된다. 따라서 단지 표면적으로 미디어에 대한 문해력 교육을 넘어서는 AI 기술의 작동 방식과 운용 원리 등 AI 기술에 관한 근본적인 이해를 바탕으로 비판적으로 활용할 수 있는 AI 리터러시 교육의 필요성이 점점 더 강조되고 있다(세계일보, 2023가).

사회적가치연구원은 2023년 5월 AI에 대한 사용자 현황 조사를 실시하였다. 전체 1,000명 중 37.6%는 AI에 대해 ‘잘 알고 있다’고 응답하였고, ‘모르고 있다’고 응답한 비율인 12.1% 대비 3배 이상 높았다. 텍스트, 이미지, 음악과 같은 새로운 콘텐츠를 만드는 ‘생성형 AI’를 사용해 본 경험이 있는 사람들의 비율은 31.1%를 나타냈고, 사용해 본 경험이 없는 사람들의 비율은 56.1%로 사용 경험 비율 대비 2배 이상 높았다. 일반인을 대상으로 AI에 관한 질문을 했을 때, AI 기술에 대해 들어보고 알고는 있지만 실제 사용해 본 경우는 많지 않은 상황이다. 조사 결과 성별, 연령별, 지역별로 AI 리터러시 격차를 나타내고 있다(사회적가치연구원, 2023).

AI의 개념은 과거에 비해 일반화되고 있으며, 사회 전반에 있어서 AI가 가져올 변화에 대한 기대와 우려는 커지고 있다(이한신, 김판수, 2019). 2023년 12월 구글은 광고 부문에 AI 기능을 도입한 여파로 3만 명에 달하는 광고 판매 부문에 대해 대규모 개편할 예정이라는 보도가 되었다(중앙일보, 2023).

생성형 AI를 기반으로 한 대화형 AI 서비스인 ChatGPT가 2022년 11월 30일 처음 세상에 공개된 이후 전세계적으로 급속한 확산이 있었다. 하지만 AI 기술을 이해하고 효과적으로 활용할 수 있는 AI 리터러시의 중요도, 그리고 그러한 신기술에 지각하고 사용하게 될 때 어떤 요인들이 영향을 미치는지에 대한 연구는 현재 충분히 이루어지지 못한 상태이다(윤종현, 윤한성, 2024). 효과적인 AI 활용은 이제 인간에게 기본적이고 필수적인 요구사항이 될 것이며 더 나은 사회를 만들기 위해서 적극적인 준비와 대응이 필요한 시점이다.

제 2 절 연구방법 및 범위

본 연구는 선행연구를 통해서 검증된 기술수용모델의 구성 변인인 지각된 유용성, 사용용이성이 AI 리터러시 수준과 AI 기술의 사용의도에 어

떠난 영향을 미쳤는지를 규명하고 지각된 유용성과 사용용이성이 사용의도 간의 경로에서 AI 리터러시 수준이 매개효과가 있는지를 파악하여 AI 기술 수용의 효과를 향상할 수 있는 방향을 제시하고자 한다.

이를 위하여 가장 먼저 국내외 선행연구를 폭넓게 검토하고 이론적인 배경이 확보된 지각된 유용성, 사용용이성과 사용의도로 구성된 기술수용모델의 개별 구성 차원과 이들의 관계에 관한 연구모형과 가설을 설정하였다.

본 연구를 수행함에 앞서 AI 기술, 리터러시와 사용의도에 대해 수행되었던 선행연구들은 다음과 같은 주제들로 진행되었다. 첫째, AI가 인간 사회에 미치는 영향에 관한 연구(김주은, 2019), 디지털 리터러시와 AI 리터러시의 개념과 관련된 탐색적 연구(이유미, 2022), AI 챗봇 발전에 따른 리터러시 필요성과 관련된 연구(이철승, 백해진, 2023) 등이 있다. 이는 AI 기술로의 패러다임 전환과 신기술 도입기 때의 기본적인 소개와 영향에 대한 탐색적 연구들이 진행되었다.

둘째, AI 기술의 사용자 성향, 위험 인식, 지각된 유용성, 지각된 사용용이성, 자기효능감 등과 같은 다양한 변수와 사용의도 간에 인과관계에 대한 연구들이 수행되었다(곽준도, 전종우, 2023; 안승규, 고상훈, 최혁진, 황순현, 광기영, 안현철, 2023; 김윤경, 2022; 고윤정, 2022).

셋째, 기술수용모델(Technology Acceptance Model: TAM)을 통한 AI 기술의 지속 사용의도에 관련 구조적 영향을 분석한 연구들이 다수 있다(장창기, 성욱준, 2022; 이한신, 김판수, 2019; 윤태환, 왕새롬, 2023).

본 연구의 연구 주제는 AI 사용 현황과 기술수용모델을 결합하여 연구모델을 구성하고 핵심 변수를 추출하였다. 본 연구는 기술수용모델의 구성 요인인 지각된 유용성, 사용용이성, 사용의도 간의 영향 관계를 기본 분석 틀로 설정하였다. 또한 이전 선행연구들에서는 실증적으로 검토되지 않았던 AI 리터러시의 성별, 연령별, 학력별, 지역별 격차 수준을 조사하고 사용의도와의 관계를 실증 분석한다.

수집된 자료는 SPSS 27.0 프로그램을 이용하여 표본의 일반적인 특성을 파악하기 위해 빈도분석과 기초 기술통계량 분석을 통해 데이터를 검

토하였다. 또한 AMOS 29.0 프로그램을 사용하여 구조방정식모델의 구성 변수들 간의 관계를 분석하였다. 확인적 요인분석을 통해 측정항목과 모형의 타당도 평가를 실시하여 기본 가설을 검정하였다. 부트스트래핑법을 통한 매개효과 검정을 거쳐, 간접효과를 확인하기 위해 경로분석을 실행하였다. 이를 통해 지각된 유용성과 사용용이성과 사용의도 간에 AI 리터러시 수준이 매개변수로서 영향을 미치는지, 그리고 그 영향을 미치는 주요 요인들이 무엇인지 확인하여 AI 리터러시 격차 해소를 위한 함의와 시사점 그리고 정책 추진 방안을 결론으로 제시하고자 한다.

제 3 절 연구의 구성

본 연구는 총 5개 장으로 구성되었고 각 장의 주요 내용은 다음과 같이 진행하고자 한다.

제 1 장 서론에서는 연구의 배경 및 목적을 기술하고 연구의 방법과 범위, 연구의 구성에 대하여 제시하였다.

제 2 장 이론적 배경에서는 본 연구의 주요 개념인 AI 정의와 최신 AI 기술 트렌드, 기술수용모델과 그 구성요인인 AI의 사용의도, 지각된 유용성과 사용용이성, AI 리터러시의 개념과 영향, AI 리터러시 활용과 격차에 대해 선행연구를 정리하여 제시하였다.

제 3 장 연구설계 부분에서는 위에서 살펴본 선행연구들을 통한 이론적 배경을 바탕으로 기술수용모델과 AI 리터러시 간의 관계에 관한 기존 연구결과를 활용하여 연구모형 및 가설을 설정하였고, 본 연구에 필요한 설문자료의 수집 및 분석 방법을 기술하였다.

제 4 장 연구 결과에서는 표본의 특성과 변수에 대한 인구통계학적 특성 및 기술통계 분석을 하고, 상관관계 분석 및 확인적 요인분석을 통해 연구모형을 검증한 후 연구가설을 검증하기 위하여 독립변수(지각된 유용성, 사용용이성), 매개변수(AI 리터러시 수준), 종속변수(사용의도) 간 영향도를 분석하고, 지각된 유용성, 사용용이성과 사용의도 사이에서 AI 리

터러시의 매개효과를 제시하였다.

제 5 장 결론 부분에서는 연구결과를 요약 정리하고, 가설검증 결과로부터 도출되는 학문적 시사점 및 정책적 시사점을 제시하고 끝으로 연구를 진행하면서 확인된 연구의 한계점과 향후 연구 방향을 제시하였다.

제 2 장 이론적 배경

제 1 절 AI 기술

1) AI 개념

AI는 4차 산업혁명 시대의 핵심으로써 전산업·사회에 파급되는 범용 기술로 미래 경쟁력을 좌우하고 혁신 성장을 위한 핵심 동력으로 발전하고 있다(이철승, 백혜진, 2023; 4차산업혁명위원회, 2021). 인공지능(Artificial Intelligence)이라는 단어는 1955년 미국 다트머스 칼리지 수학 교수 John McCarthy가 처음 사용하였다. 그는 MIT와 벨연구소, IBM에 있던 동료들과 함께 “언어를 사용하고 추상적인 개념을 만들며 인간이 풀게 되어 있는 온갖 종류의 문제를 풀고 스스로 개선할 수 있는 기계를 만드는 방법을 찾아내는 작업”에 착수했다(Anthes, 2017). AI는 인간의 지적 능력을 기계로 구현하는 과학기술을 말한다(Elsholz, Chamberlain, & Kruschwitz, 2019). 딥러닝 기술을 학습하여 최적의 답을 찾아내고, 추론 및 예측을 하며, 주어진 문제를 스스로 발견하고 해결하는 행동 단계에 이르기까지 다양한 분야의 연구가 활발히 진행되고 있다(이철승, 백혜진, 2022; 조민호, 2021). 또한 ‘AI’ 용어의 범위는 일반적으로 인간의 지능이 필요한 작업을 수행하는 기술이나 시스템을 포괄하는 광범위한 분야에 퍼져 있으며, 자연어처리, 이미지 인식, 의사결정, 문제해결 등이 모두 포함된다.

2023년 3월 EU 대표들은 AI를 ‘다양한 수준의 자율성을 지닌 채, 명시적 또는 암시적 목표를 위해 물리적 또는 가상 환경에 미치는 영향을 예측하고 권장 사항이나 의사결정을 출력하도록 설계된 기계 기반 시스템’이라고 정의했다(Euractiv, 2023). 이 정의에서 중요한 것은 AI가 ‘권장 사항’을 추월한 ‘의사결정’을 출력한다는 것이며, 이는 AI가 여러 환경을 고려한 최적의 선택을 권장해 줌으로써 인류는 많은 이점을 얻을 수 있다. 그러나 다른 한편으로 예상하지 못했던 위협과 도전도 동시에 가져올 수

있다는 점에서 양날의 검이라는 긍정과 부정의 양면성으로 표현될 수 있다.

2) AI 발전의 역사

2022년 ChatGPT 출시로 인해 AI의 역사는 다시 한번 획기적인 전환기를 맞게 되었다. 1955년 다트머스 회의에서 ‘인공지능’이라는 용어가 처음 사용된 이후 현재까지 몇 차례 변곡점이 있었다(박태웅, 2023).

최근 AI 붐이 일어나면서 이전에 없었던 신기술을 접하고 있는 듯하지만, 사실 지금은 제3의 부흥기라고 할 수 있다. 제1의 부흥기는 1960년대와 1970년대 초가 해당한다. 이때에는 AI 이론에 대한 실행이 뒷받침되었지만, AI 기술이 약속을 이행하지 못하면서 1차 AI 겨울을 맞이하게 된다. 당시 연구자들은 AI의 잠재력에 대해 낙관적으로 생각했으나 그들이 사용하던 알고리즘과 모델이 현실 사회의 복잡한 문제를 해결할 수 없다는 사실을 깨닫게 된다. 이에 따라 AI 연구에 대한 투자가 위축되고 기술에 대한 신뢰가 떨어지게 되었다.

제2의 부흥기는 1980년대 말과 1990년대 초에 찾아왔다. 이때의 부흥은 전문가 시스템의 구축으로 시작되었다. 그러나 해당 시스템의 한계가 드러나면서 두 번째 겨울이 찾아오게 된다. 전문가 시스템은 규칙과 논리에 의존하여 의사결정을 내리는 AI 기술이다. 그러나 이 시스템은 불확실성을 처리하거나 데이터를 통한 학습을 할 수 없었다. 결과적으로 AI 연구에 대한 투자와 지원은 다시 감소했고 많은 연구자가 이 분야를 떠날 수밖에 없었다.

현재는 제3의 부흥기라고 할 수 있으며, 머신러닝과 딥러닝 기술의 발전, 양질의 빅데이터 가용성, 하드웨어와 클라우드 컴퓨팅 성능 향상에 힘입어 최근 몇 년 동안 AI는 급속도로 발전하게 되었다. 이제 모든 분야에서 AI가 언급되고 있으며, 생활 전반에 영향을 미칠 것으로 예상된다.

3) AI 발전단계

AI가 빠르게 발전하면서 다양한 기술이 인간 삶에 영향을 미치고 있다.

전문가들은 AI를 지식과 자율성 수준에 따라 강한 AI(Strong AI, 일반적인 인공지능)와 약한 AI(Weak AI, 좁은 의미의 AI)의 두 가지 유형으로 분류한다(Morris et al., 2023). 강한 AI는 여러 영역에서 인간처럼 지능적으로 작동하도록 설계되는 것이며, 약한 AI는 특정 작업 또는 일련의 작업을 수행하도록 설계되는 것이다.

최근 OpenAI는 AI의 발전단계를 5단계로 나누어 각 단계가 인류에게 어떤 의미를 가지는지, 그리고 AGI(Artificial General Intelligence, 인공일반지능)와 어떤 관계가 있는지 설명하고 있다. OpenAI가 제시한 AI 발전 과정 5단계는 단순한 도구로서의 AI에서 시작해 궁극적으로는 인간과 동등하거나 인간보다 뛰어난 수준으로 발전할 가능성을 염두에 둔 것이다(Bloomberg, 2024).

다섯 단계는 궁극적으로 AGI에 도달하는 과정을 나타낸다. AGI는 인간처럼 다양한 작업을 수행할 수 있는 AI를 의미하며 인간의 지능과 동등하거나 이를 넘어서는 수준을 목표로 한다. 1단계에서 5단계로 갈수록 AI는 점점 더 인간과 비슷한 능력을 갖춰 결국 AGI에 도달할 것이다. 1단계는 이미 현실화하여 대화형 챗봇은 다양한 산업에서 사용 중이다. 2단계도 이미 구현되어 다양한 문제해결에 도움을 주고 있다. 3단계는 자율주행 자동차나 드론과 같은 기술을 통해 실현되고 있으며 가까운 미래에는 완전히 실현되어 더 많은 곳에서 사용될 것이다. 4단계는 연구소와 기업에서 활용되기 시작했지만, 아직 널리 보급되지 않았다. 마지막 5단계는 현재 연구 단계로 현실화할 때까지 수십 년 이상 걸릴 것이다.

구글의 딥마인드(DeepMind) 연구진도 인공일반지능(Artificial General Intelligence, AGI)을 6단계로 구분하였다. 딥마인드는 특히 인간처럼 사고하고 학습하는 AI를 개발하는 데 집중하고 있다. 이런 차원에서 AGI로 가는 과정을 체계적으로 설명하는 단계들을 제시한다. 딥마인드는 AI가 단순한 기술이 아닌 인간의 지능을 증강하거나 심지어 대체할 수 있는 중요한 도구가 될 수 있다고 본다. ChatGPT, 바드, 라마2 등과 같은 생성형 GPT를 1단계인 유망한(Emerging) AGI로 평가했다. 아직 ChatGPT가 시작 단계이며, 앞으로 인간의 능력을 100% 발휘하거나 뛰어넘을 수 있는 AI가 개발될 수 있음을

내포하는 분류이다(Morris et al., 2023).

딥마인드는 AI가 어떻게 점진적으로 발전할 수 있을지 그리고 그 발전이 인류에게 어떤 영향을 미칠지에 대한 청사진을 제공한다. 딥마인드 역시 AGI 도달이 목표이며 AI의 각 발전단계에서 어떤 역할을 할지에 집중한다. 딥마인드가 제시한 AI의 각 단계는 결국 AGI로 향해 가는 과정이다. 첫 번째 단계인 제한된 AI에서 시작해, 점점 더 인간과 비슷한 능력을 갖춘 AI로 발전해 나가는 것이 목표이다. 궁극적으로는 AGI에 도달하여 인간과 같은 수준, 혹은 그 이상의 지능을 갖춘 AI가 될 수 있다.

1단계는 이미 구현되어 있으며 알파고와 같은 AI가 그 예이다. 2단계는 범용 AI로 가기 위한 연구가 진행 중이지만, 아직 완전한 범용 AI는 존재하지 않는다. 이 단계의 AI는 수십 년 내에 실현될 가능성이 있다. 3단계는 자율 학습을 통해 다양한 작업을 처리할 수 있는 AI로 알파제로 같은 예가 있으며, 더 많은 응용 가능성이 연구되고 있다. 4단계인 인간 수준의 AI는 아직 이론적인 단계에 머물러 있다. 이 단계에 도달하기까지는 상당한 시간이 필요할 것이다. 5단계인 AGI는 현재 연구의 최종 목표로 언제 도달할 수 있을지는 불확실하지만, 수십 년 내 혹은 그 이상의 시간이 필요할 것이다.

OpenAI는 초기 단계에서부터 구체적인 대화형 AI에 초점을 맞추며 점차 창의적이고 독립적인 행동을 할 수 있는 AI로 발전하는 것을 목표로 한다. 반면 구글 딥마인드는 제한된 분야에서 시작해 점차 범용적인 AI, 그리고 자율 학습과 적응을 통해 인간 수준의 지능에 도달하는 방향을 추구한다. 두 회사 모두 AGI를 목표로 하고 있지만, OpenAI는 조직의 역할을 할 수 있는 AI로의 발전을 강조하는 반면, 구글 딥마인드는 인간 수준의 범용적인 AI를 구현하는 데 중점을 두고 있다.

한편 챗봇(Chatbot)과 ChatGPT는 모두 AI 기술을 사용하는 대화형 소프트웨어이다. 챗봇은 미리 정해진 데이터베이스에서 사용자의 질문에 대한 답변을 제공하는 반면, ChatGPT는 대형언어모델(Large Language Model: LLM) 기반의 정교한 처리 및 조정 엔진을 사용하여 실시간으로 코드를 제안하고 코드 라인을 완성함으로써 개발자가 코드를 더욱 효율적으로 작성할 수 있도록 설계되었다(장창기, 성욱준, 2022). 본 논문에서는

AI와 AI 기술을 광의의 개념으로, ChatGPT 등과 같은 생성형 AI는 협의의 개념으로 사용코자 한다.

제 2 절 기술수용모델

첨단 기술 및 IT 신제품의 수용 연구에서 많이 응용되고 있는 기술수용모델(TAM: Technology Acceptance Model)은 기술 수용자의 사용 행동과 사용의도를 설명하기 위한 연구모형이다(Davis, Bagozzi, & Warshaw, 1989). AI와 같은 새로운 기술 이용자의 수용 태도, 지각된 유용성, 사용용의성, 그 영향력과 사용의도간의 관계를 분석하는 데 합리적인 모델이다. 기술수용모델은 특정 혁신에 대해 구성원들이 가지고 있는 긍정적인거나 부정적인 태도, 믿음, 사용의도와 실제 사용 간에 어떠한 인과관계가 있는지와 수용 과정에 영향을 미치는 외부 요인들이 무엇인지 발견하는 데 초점을 두고 있다(Davis, Bagozzi, & Warshaw, 1989).

기술수용모델은 태도를 통해 행위를 예측하는 행동 의도 모델인 합리적 행동이론(TRA: Theory of Reasoned Action)을 이론적으로 확장하여 도입하였다(Ajzen & Fishbein, 1975, 1980). 합리적 행동이론은 태도, 신념, 의도와 행동 등의 관계를 검증하는 모델이며, 기술수용모델은 기술을 채택하는 데 있어 지각된 사용용의성, 지각된 유용성, 이용태도, 사용의도와 행동 간의 관계 검증이 필요하다(Davis, 1989). 기술수용모델은 사용자가 기술을 수용할 때 어떤 요인들이 결정적 영향을 미치는지를 설명하는 것이며 사용자의 태도와 연관되며 그 태도는 행위의도, 사용의도에 영향을 미치며 그 의도는 실제 행위에 영향을 미친다는 가정에 따라 사용자 수용은 지각된 사용용의성과 지각된 유용성에 의해 이루어진다(김태문, 한진수, 2009).

기술수용모델의 최종 목적은 내부적으로 형성된 태도와 의도에 대해 외부요소가 미치는 영향의 비중을 반영하여 새로운 기술이나 제도의 수용에 대한 결정요인을 포함하는 사용자의 행위를 설명하는 데 있다(남일규, 김승권, 김준연, 김영중, 박희준, 2017). 지각된 유용성은 한 개인이 특정한 시스템을

사용하는 것이 자신의 직무 성과를 향상할 수 있다고 믿는 정도로 정의하며 사용용이성은 한 개인이 특정한 시스템을 사용할 때 신체적 정신적 노력에서 자유롭다고 믿는 정도로 정의된다(Davis, Bagozzi, & Warshaw, 1989). 기술 수용모델에서 지각된 유용성(Perceived Usefulness)과 지각된 사용용이성(Perceived Ease of Use)의 두 개념을 이용하여 새로운 혁신 기술을 수용하려는 태도와 행동 의도 간의 관계를 검증하였다(Davis, Bagozzi, & Warshaw, 1989).

기술수용모델에 관한 연구들이 지속적으로 진행되면서 지각된 유용성과 사용용이성 간의 다양한 연구들이 수행되었다. 이를 기술 분야별로 살펴보면 모바일 인터넷, 사물인터넷, 챗봇과 같은 IT 관련 기술에 대한 소비자의 수용 의도에 관한 연구(김윤경, 2022; 조진완, 이종호, 2014; 오종철, 윤성준, 우원, 2010), 혁신적 관점에서 혁신 정보기술에 대한 소비자의 수용도를 측정한 연구(Shin, 2009), AI 기술 수용의도에 관련한 연구 등으로 구분할 수 있다(장창기, 성욱준, 2022; 곽준도, 전종우, 2023).

기술수용모델은 소비자의 신념, 태도 그리고 행동 의도에 대한 외부 요소들의 영향력을 추적할 수 있고 정보기술 및 제품들의 수용과 확산을 설명하는 간결하면서 강력한 특성이 있다는 점에서 매우 큰 의의가 있다고 할 수 있다(Davis, 1989; Venkatesh & Davis, 1996). 기술수용모델은 다양한 과학기술과 제도의 특성에 따라 확장된 기술수용모델로 응용되었다(고영화, 임춘성, 2021). 신기술의 특징에 따라 이용 의도에 영향을 미치는 외부 변수 또한 달라질 수 있다(장한진, 노기영, 2017). 따라서 새로운 기술에 대한 이용 의도는 지각된 유용성과 사용용이성에 의해 결정되거나 특정 기술에 특화된 외부 요인들이 유용성과 사용용이성에 영향을 미친다는 사실이 기존 모델에 반영되었다(Venkatesh & Morris, 2000; 심태용, 윤성준, 2020).

제 3 절 사용의도

사용의도(Intention to Use)는 소비자가 특정 제품을 구매하고 사용하는 것에 대한 의사를 나타내는 표현이다. 특정 제품이 소비자에게 주는 주관적 경험을 연구하거나 시장의 수용도를 평가할 때 항상 사용된다. 계획된 행동이론(Theory of Planned Behavior)에서 행동 의도는 개인의 태도, 주관적 규범 및 지각적 행동 통제의 조합을 기반으로 형성된다(Ajzen, 1991). Davis(1989)가 제안한 기술수용모델은 기술에 대한 사용자의 지각된 유용성(Perceived Usefulness) 및 지각된 용이성(Perceived Ease of Use)이 사용 태도에 영향을 미치고 사용 의도에도 영향을 미치고 마지막으로 사용 행동에 영향을 미친다(Davis, Bagozzi, & Warshaw, 1989).

AI의 사용의도에 대한 연구에서는 영향 요인을 기술적 요인과 사회적 요인으로 나눌 수 있다. 첫째, 기술적 요인은 지각된 유용성과 지각된 사용용이성이 포함된다. 유용성과 용이성을 감지하는 것이 차별화된 추천 서비스에 긍정적인 영향을 미친다는 연구가 있다(Kim & Nam, 2019). 둘째, 사회적 요인은 사회적 영향, 개인적 특성, 문화적 배경에 따라 결정되는 것이다. 선행연구에 따르면 사용자가 공리적인 목적으로 AI 서비스를 선택하기 쉽다(Wang & Hou, 2019).

새로운 기술에 대한 개인의 태도는 그 기술에 대한 신뢰와 내재적 관여(Intrinsic Involvement)를 통해 형성되는 개인의 기술에 대한 인식과 관련이 있다(장창기, 성욱준, 2022). 새로운 기술에 관한 신뢰는 지각된 위험을 해소하는 선행요인이 되고, 이용자로서 기술의 발전 과정에서의 관여는 새로운 기술을 도입하는 데 긍정적인 영향을 미친다(Jackson, Chow, & Leitch, 1997; Pavlou, 2003). 특히, AI 기술에 대한 편리성에 관한 인식은 이용자와 공급자 모두에게 긍정적으로 받아들여지고 있다(윤상오, 2018). AI 기술에 대한 부정적 영향보다는 개인이 인식하는 혜택에 관한 긍정적 인식이 더 큰 영향을 나타내고 있다(장창기, 성욱준, 2022).

제 4 절 지각된 유용성

지각된 유용성(Perceived Usefulness)이란 “어떤 사람이 특정한 시스템을 이용하여 일의 성과를 높일 수 있다고 믿는 정도”를 말한다. 여기서 유용하다는 것의 의미는 특정한 기술이나 서비스 등을 선택했을 때 이득을 얻거나 조직 내 상황의 경우 보상을 받거나 좋은 성과를 강화해 준다는 것과 같이 기술을 수용했을 때 얻을 수 있는 혜택을 의미한다고 볼 수 있다(성동규, 2009). 지각된 유용성이라는 개념에 대해 정보 시스템 연구 및 신기술 도입 시기에 대표적으로 많이 거론되고 있는 기술수용모델은 1989년 Davis에 의해 처음 소개되었으며, 아주 견고한 모델로 평가받으며 정보 시스템과 정보화기기 등에 관련된 많은 연구에 적용되었고, 지각된 유용성은 사용자들의 수용 전과 수용 후의 인지 정도 및 신념을 검증하도록 하고 있다(성동규, 2009).

온라인 서비스에 관한 지각된 유용성은 수용 후에도 대상 서비스의 지속 사용 의지와 사용자 만족에 유의한 영향을 준다는 것을 검증하였다(Bhattacharjee, 2001). AI 기술 유행 이전 많이 거론되었던 스마트폰과 소셜 네트워크 서비스에서도 지각된 유용성은 대상 정보기술 서비스 사용자들의 지속 사용 의지 형성에 중요한 역할을 담당한다고 확인된 바 있다(임성진, 2012). 즉, 어떤 한 제품의 수명을 지속해서 존속시키기 위해서는 신제품이 기존의 제품보다 성능이나 기능 측면에서 전달해 줄 수 없었던 가치를 고객에게 제공할 때 유용성이 높게 측정되며, 시장에 빠르게 수용될 수 있다는 것이다(Rogers, 2003).

기술수용모델이 기술 채택이나 이용 의사가 지각된 유용성과 지각된 사용용이성에 근거한다고 전제하여 연구를 수행한 바 있다(성동규, 2009). 온라인 서비스에 관한 지각된 유용성은 수용 후에도 대상 서비스의 지속 사용 의지와 사용자 만족에 유의한 영향을 준다는 것을 검증하고 있었다(Bhattacharjee, 2001). 기술수용모델은 원래 정보기술의 채택을 설명하면서 출발했지만 이후 정보기술 사용을 넘어 다양한 부분에서 활용되고 있다(Venkatesh, 2000). 이처럼 기술수용모델을 활용한 연구들이 다양

하게 수행될 수 있었던 이유는 기술수용모델이 신기술의 등장과 이를 사용하는 수용자의 태도와 행위를 설명하는 데 유용하기 때문이다.

기술수용모델은 합리적 행동이론의 영향으로부터 시작되었다. 합리적 행동이론은 인간의 행동 변화를 설명하는 이론 가운데 가장 광범위하게 이용되고 있는 이론으로서 합리적인 이성에 근거한 주관적 판단이 수용자의 행동을 결정한다고 본다(임성진, 한경석, 정미라, 2017). 합리적 행위 이론과 마찬가지로 인간의 행위에 주목하는 기술수용모델은 설명력이 높고, 간결한 모델을 제시할 수 있다는 장점이 있다고 평가되었다(Venkatesh, 2000). 기술수용모델은 처음 소개된 이후 아주 견고한 모델로 평가받으며 많은 정보시스템 연구에 적용되었으며, 지각된 사용용이성, 지각된 유용성으로 구성된 내부 신념, 태도, 의도에 대한 외부 변수의 경로를 확인할 수 있는 매우 설명력이 높은 모델이다(Davis, 1989). 내부 신념인 지각된 사용용의성과 지각된 유용성의 정의를 살펴보면, 지각된 사용용의성은 특정 기술을 사용하는 데 있어 적은 노력이 필요할 것이라는 믿음의 정도 그리고 지각된 유용성은 특정 기술의 사용이 자신의 직무 성과를 높여 줄 것이라는 믿음의 정도로 정의된다(Davis, 1989).

제 5 절 사용용이성

지각된 사용용이성(Perceived Ease of Use)은 정보시스템 인터페이스에 대한 사용자의 평가로서 입력 및 출력의 용이성, 검색 및 분석 과정의 용이성, 도움말 기능의 다양성과 편리성 등으로 사람이 기술을 이용할 때 노력을 들이지 않는 정도를 말한다(임성진, 2012).

Davis(1989)는 지각된 사용용이성(Perceived Ease of Use)을 “잠재적 이용자가 특정한 정보기술 혹은 시스템을 이용하는 것이 신체적이거나 정신적 수고가 적게 들 것이라고 믿는 정도” 또는 “잠재적 사용자가 큰 노력 없이 새로운 기술을 사용할 수 있을 것으로 기대하는 정도”라고 정의하면서 사용이 편리한 기술은 상대적으로 편리하지 않은 기술보다 이용자들이

활용하는 비율이 높다는 연구 결과가 많이 제시되고 있다. 기술수용모델을 구성하는 주요 개념인 지각된 사용용이성과 지각된 유용성 간의 관계에 대해 기존의 다양한 선행연구들에서 사용용이성은 지각된 유용성에 유의미한 영향을 미쳤으며, 지각된 사용용이성이 지각된 유용성의 선행변수임을 보여주고 있다. 이는 사용이 쉬운 시스템이나 기술은 그렇지 않은 것보다 더 잘 사용할 수 있고 업무 실행 효과도 더욱 높아진다는 것을 의미한다. 또한 지각된 사용용이성이 사용의도에 직접적인 영향을 미친다는 것은 사용자의 수용 정도를 직접적으로 향상할 수 있다는 것을 의미한다(Davis, 1989).

사용자가 제품의 사용법을 습득하는 정도가 빠를수록 신제품이 시장에서 수용되는 속도가 빠르다는 사실이 확인되었다(Rogers, 2003). 또한 이용자가 느끼는 제품 이용의 어려운 정도가 실제로 어떤 서비스를 선택할지를 결정하는데 높은 상관관계가 있음이 확인되었다(Ajzen, 1991). Davis(1989)의 기술수용모델과 Delone과 McLean(1992)의 정보시스템 성공모델을 통합하여 연구한 Rai, Lang, & Welker(2002) 역시 사용용이성이 시스템 사용과 관련된 태도에 큰 영향을 미치며 나아가 사용자 만족에 영향을 줄 뿐만 아니라 시스템에 대한 개인적 가치평가인 개인적 영향(Individual Impact)과도 인과관계가 있다고 밝혔다(Rai, Lang, & Welker, 2002).

기술수용모델은 컴퓨터 기술 수용 행위를 설명하기 위해 합리적 행동이론을 수정하여 제시하였다. 본 모델은 두 개의 신념, 즉 지각된 사용용이성과 지각된 유용성에 의해 기술 수용 여부를 결정하는 것이다. 두 가지 핵심 믿음인 지각된 유용성과 지각된 사용용이성이 태도와 행동 의도를 매개로 하여 실질적으로 이루어진다. 이러한 연구를 통해 실질적인 시스템 사용 행위와 의도 사이에는 서로 높은 상관관계를 유지하고 있음이 밝혀졌으며, 의도는 태도에 의해 결정된다는 것이 검증되고 있다(Venkatesh, 2000). 이렇듯 기술수용모형은 현재 첨단 기술 연구 분야에 있어서 기술 수용과 관련한 수용자의 인지적 특성에 관한 연구로써 그 타당도를 널리 인정받고 있는 이론이라 할 수 있다(임성진, 2012).

제 6 절 AI 리터러시

1) AI 리터러시의 개념

리터러시(literacy)라는 용어는 단순한 ‘문해력’의 의미를 뛰어넘어 ‘소양’이라는 의미로 쓰임이 확장되었다(이철승, 백혜진, 2022). 다양한 단어와 결합하여 사용되고 있으며 대표적으로 미디어 리터러시와 디지털 리터러시가 있다(이우진, 백혜진, 2022). 미디어 리터러시는 텔레비전과 같은 영상 장치가 등장하고 방송이 실현되면서 본격적으로 개념으로 자리 잡았다(이철승, 백혜진, 2023). 디지털 리터러시는 인터넷과 모바일 디바이스의 사용 및 소셜 미디어의 확장으로 등장하였으며, 단순히 기기를 사용하는 방법을 넘어 정보를 다루고 가공하는 일까지 의미하게 되었다(Jenkins, 2009). 이후 가상비서, 챗봇, 안면인식 기술, 자율주행 자동차에 이르기까지, AI 기술이 급속도로 발전하고 광범위하게 보급됨에 따라 AI 리터러시의 개념이 등장하였다. 이러한 다양한 리터러시는 각 분야에 대한 기본적인 소양뿐만 아니라, 비판적 사고를 바탕으로 분석하고 활용할 수 있는 능력을 말한다(이철승, 백혜진, 2023).

현재 우리는 PC, 모바일과 네트워크를 넘어 모든 사물이 인터넷으로 연결된 초연결 시대에 살고 있다. 모든 대상과 네트워크 연결이 가능해짐에 따라 상호 교환되는 데이터양과 정보량은 이전 지식사회 대비 비교할 수 없을 정도로 기하급수적으로 증가하고 있다. 이러한 정보의 바다와도 같은 사회에서 우리는 공유된 정보를 통합적 관점에서 비판적으로 수용하고 활용할 수 있는 능력인 디지털 리터러시가 요구된다(고운정, 2022). 디지털 리터러시는 읽고 쓰고 이해하는 문해력이라는 의미에서 진화하여 시각 리터러시, 텔레비전 리터러시, 컴퓨터 리터러시, 멀티미디어 리터러시, 정보 리터러시, 정보통신 리터러시, 미디어 리터러시, 디지털 리터러시, AI 리터러시, 데이터 리터러시 등으로 확장되었다(Emwanta & Nwalo, 2013; 권성호, 김성미, 2011). 지금까지 리터러시가 진화되어 온 과정을 [표 2-1]과 같이 정리하였다.

[표 2-1] 리터러시 개념 변화¹⁾

리터러시 종류	시기	주요 특징
시각 리터러시	1960년대	이미지로부터 의미를 구성하는 능력
텔레비전 리터러시	1950년대 이후	TV 보편화 이후 시각적 지각 능력과 비판적으로 시청하는 능력
멀티미디어 리터러시	1990년대	매체 내에서 그림, 사진, 동영상 등 멀티미디어 기능을 활용할 수 있는 능력
정보 리터러시	1990년대 이후	인터넷 확산 이후 정보를 획득해 이용 및 평가할 수 있는 능력
미디어 리터러시	1990년대 이후	정보, 컴퓨터, 영화 및 비디오 리터러시 등의 개념을 포괄한 메시지를 해석하는 능력
디지털 리터러시	2000년대 이후	기술과 기능 활용을 초월한 문제해결 과정에서 필요한 정보 인식과 비판적 평가와 같은 능력
AI 리터러시	2015년 이후	AI 기술에 대한 이해와 인공지능과 효과적으로 소통하고 협업, 비판적으로 수용할 수 있는 능력

디지털 리터러시 연구는 초기 디지털 리터러시의 개념 및 구성요소에 초점을 맞추고 있으며, 점차 사회·문화·윤리·소통 영역으로 확대되었다(한정선, 오정숙, 2006; 권성호, 김성미, 2011). 최근에는 다양한 분야의 융·복합적·창의적 인재 양성에 초점을 두고 있는 추세다(신소영, 이승희, 2019).

AI 리터러시 개념도 디지털 리터러시에서 세분화하여 발전한 것으로 현재 AI 리터러시에 대한 개념 정립에 관한 연구는 시작 단계이며 다양한 연구 문제 검증을 위한 구성요소와 도구를 개발하고 확립하는데 중점을 두고 있으나 미흡한 실정이다(고운정, 2022). AI 리터러시의 개념 정립에 관하여 Long과 Magerko(2020)는 AI 리터러시를 AI의 협업 및 소통 기술을 집이나 직장에서 도구로 사용하고 평가할 수 있도록 하는 능력이라고 정의하였으며, AI의 개념, AI의 능력 및 AI 작업 방법, 활용 방법 및

1) 황현정, 박지수, 김승완. (2023). “국내외 AI 리터러시의 연구 동향 분석: 체계적 문헌 고찰 방법을 중심으로.” 『교육문화연구』, 29(3), 135-164.

인식에 대한 방법으로 구성요소를 제시 하였다(Long & Magerko, 2020). 한편 AI 리터러시가 디지털 사회에 참여할 수 있도록 하는 새로운 능력이나 방법이라고 설명하였으며, AI 역량을 갖추기 위해 AI에 대한 자기효능감, 의미 부여, 창조적 자기효능감의 중요성을 제시하였다(Kong, 2008).

국내에서는 Long과 Margerko(2020)의 연구를 바탕으로 AI 리터러시를 개인이 AI 기술에 대한 비판적인 평가능력, AI와 효과적인 소통 및 협업 능력, AI의 기본개념 및 원리 이해, AI 기술을 활용한 문제해결 및 도출 능력(김태령, 류미영, 한선관, 2020; 이철현, 2020; Ruy & Jo, 2021) 등으로 정의한다. AI 리터러시 정의에 기반하여 초·중등학교에서 각 과목에서 다양한 교육 프로그램의 개발 및 활용을 통해 AI 리터러시가 향상되고 있는지 확인하는 연구가 대부분이다(이다겸, 김성원, 이영준, 2021; 김성주, 2021; Kong, 2008; 김진석, 2021). 그 이외 AI 윤리를 다루고 있는 연구도 드물게 찾아볼 수 있다(김태창, 변순용, 2021; 송선영, 2021). 한편 AI 기반 공공서비스 정책 수용 의도에 관한 연구에서 AI 기술에 대한 인식과 디지털 리터러시가 사용의도에 미치는 영향을 실증분석하는 연구(장창기, 성옥준, 2022)가 있다.

AI 리터러시는 AI 기술과 각종 응용 프로그램을 이해, 활용하는 데 필요한 지식, 기술 및 윤리적 인식을 포함하는 포괄적 개념이다. 이는 기술적 이해를 넘어 비판적 사고, 윤리적 고려 및 AI의 광범위한 사회적 영향에 대한 인식을 포함하는 다차원적 개념이다(최재현, 2023).

2) AI 리터러시의 영향

사회적가치연구원에서는 2023년 5월 AI 기술이 해결할 수 있는 사회 문제로서 개선되거나 악화할 것으로 예상되는 문제들을 조사하였다. 설문 결과를 보면 ‘삶의 질 저하(49.5%)’, ‘급격한 사회구조 변화(48.8%)’, ‘교육 불평등(44.6%)’ 문제를 AI 기술이 해결할 것이라는 답변이 거의 50%에 육박한 수준으로 나타났다. 특히 개선이 예상되는 사회문제로 AI 활용 분야 중 ‘교육 불평등’이 1순위로 선택되어 다른 문제 대비 평균 10% 포

인트 높은 19.2%를 나타내어 가장 기대감이 높았다.

반면, 전체의 75.3%의 국민은 ‘고용 및 노동 불안정’ 문제가 AI 기술로 인해 가장 악화할 것으로 예상하며 일자리 위협을 가장 걱정하고 있으며 국민 3명 중 1명 이상이 1순위(35.5%) 문제로 선택하였다(사회적가치연구원, 2023). 2순위는 ‘급격한 사회구조 변화’ 문제가 54.7%, 3순위는 ‘사회통합 저해’ 문제가 54.5%로 그 뒤를 이었다. 특히 ‘급격한 사회구조 변화’ 문제는 AI로 인해 악화할 문제이자 해결할 수 있는 문제로 각각 2위로 꼽혀, AI에 대한 이해 및 활용 정도에 따라 우리의 삶에 미칠 영향이 달라질 수 있음을 나타내고 있다(사회적가치연구원, 2023).

AI 기술의 빠른 발전은 우리 사회의 많은 문제들을 개선하는 원동력이 되기도 하고 때로는 새로운 격차, 고립, 차별과 침해 등 다양한 사회문제들의 시작점이 되거나 사회 시스템에 악영향을 미칠 잠재력도 보유하고 있다(이은창, 2023). 디지털 미디어 리터러시는 AI 리터러시의 포괄적 개념이라고 할 수 있으며, 디지털 미디어 리터러시 격차의 요인을 세대와 경제 수준을 중심으로 분석한 연구가 있다(안정임, 서운경, 2014). AI 기술을 업무나 실생활에서 효과적으로 활용하는 것은 이제 필수적인 요소가 될 뿐만 아니라 이를 넘어 더 나은 미래 준비를 위한 대응이 필요한 시점이다. AI 기술이 초래할 수 있는 성별, 연령별, 학력별 및 지역별 AI 기술 이해와 활용 역량 격차를 축소하기 위해 AI의 작동 방식과 기능을 이해하고 활용하는 역량을 키우고 AI가 창출하는 결과물에 대해 비판적으로 사고할 수 있는 AI 리터러시 역량 향상이 필요하다.

3) AI 리터러시 활용

AI 기술이 대중적으로 활성화되는 시점에 디지털 정보화 활용 수준 사례를 참고하여 AI 리터러시 활용 수준과 관련한 영향력을 분석할 수 있다. 디지털 정보화 활용 능력의 경우 초기에는 디지털기기 관련 보유 및 사용에 대한 ‘접근 가능성’에서 시작되었으나, 현재 디지털기기 보급과 인터넷 접근성이 포화기에 접어들면서 단순히 물리적 접근에만 국한된 것이

아닌 정보통신 기술을 ‘이용 및 활용’하지 못하는 사람 사이의 격차까지도 포함하는 광범위한 개념으로 진화되고 있다(문영임, 이성규, 김지혜, 2021). 즉 기기를 가지고 있음에도 불구하고 활용 능력이 현저하게 낮아 활용도가 떨어지는 경우가 포함되는 것을 의미하며 이 경우 정보 접근성을 기본으로 개인의 정보 활용 수준의 관점에서 디지털정보격차가 논의된다(주경희, 김동심, 김주현, 2018). AI 사용이 더욱 보편화될수록 AI 리터러시 격차가 심화될 수 있음을 시사하고 있다.

Molnar(2002)에 의하면 정보기술 적응 시기에 따라 정보격차는 세 단계로 구분되고 단계별로 정보격차 수준과 특성이 다르다. 먼저 도입기(Early Adaption)에서는 정보에 대한 ‘접근 여부’에 따라 격차가 발생하며, 이후 도약기(Take-off)에서는 일정 부분 접근 격차가 해소된 후에도 ‘사용 여부’에 따라 여전히 격차가 존재한다고 설명하였다(Molnar, 2002). 이때에는 동일한 기술과 기기를 사용하지만, 이용량에 따라 개인별 격차가 발생한다. 마지막으로 포화기(Saturation)에 이르면 정보에 대한 다양한 접근 경로와 저렴한 접근 비용으로 인해 정보 활용에 있어 ‘사용의 질’ 격차가 발생한다(Molnar, 2002). 이는 양적인 이용 격차와는 구분되는 것으로 이 단계에서는 자신이 필요로 하는 정보를 선택하거나 걸러내는 활용 능력이 중요한 요인으로 작용한다(Hargittai & Hinnant, 2008; 김효정, 2018; 문영임, 이성규, 김지혜, 2021).

정보 활용의 질적 차이에서 기인한 디지털격차가 사회적 쟁점으로 주목받는 이유는 그 파급효과가 다른 격차보다 더욱 크기 때문이다(송효진, 2014). 사회·정치·문화·경제 등 기존의 다양한 영역에서 발생하던 사회적 불평등에 더해 AI 리터러시 격차가 더해지면서 집단 내·외에서의 소외와 배제가 더욱 심화할 수 있다. 또한 AI 기술 활용 수준의 차이로 인해 격차 해소 기회조차 얻지 못하는 AI 사각지대에 속한 집단의 경우 소외와 배제의 악순환이 되풀이될 수 있다(문영임, 이성규, 김지혜, 2021).

정보를 수용하는 정도의 차이가 곧 정보의 격차로 나타나기 때문인데 정보수준에 대한 정의는 주로 세 가지 유형으로 분류되며 접근수준을 제외하고 학자별로 용어의 차이를 보이지만 내용에서 유사한 맥락을 보이며

크게 정보접근(1유형), 정보역량(2유형), 정보활용(3유형)의 세 가지 유형으로 나뉜다고 할 수 있다(오설미, 최송식, 2021). AI 리터러시 격차에도 이 이론을 적용하면 AI 리터러시 접근, 역량, 활용의 세 가지 요소가 포함된다고 볼 수 있다. 본 연구에서는 AI 리터러시를 접근과 활용 측면에 초점을 맞추어 다루고자 한다.

AI 기술 수용에 관한 기존 연구는 주로 노년공학이나 돌봄의 관점에서 노인의 건강이나 생활을 돕기 위한 연구가 대부분이었다. 2016년 이후 정보통신기술과 4차 산업혁명 시대에 관한 관심이 증가하면서 보건복지의 영역에서 AI 등과 같은 과학기술의 연구들이 증가하고 있으며 정부와 지자체, 기업에서도 노인돌봄을 위한 AI 스피커 및 IoT(사물인터넷 : Internet of Things) 등과 같은 서비스를 개발하고 있다(백민소, 신준섭, 신유선, 2021). 이를 위해 노인의 AI 기술개발에 관한 연구도 지속적으로 수행되고 있다(이기호, 2017; 이지연, 이재익, 전지호, 2017). 최근에는 AI 기술에 대한 인식과 디지털 리터러시가 지각된 유용성을 매개변수로 사용할 때 사용의도에 미치는 영향은 긍정적이며 지각된 사용용이성보다 상대적으로 큰 것으로 확인되었다(장창기, 성욱준, 2022).

AI에 관한 연구는 외부 환경에 관한 지각을 통해 자율적으로 행동하는 기계 또는 시스템이라고 하는 인간의 대리자(Agent)에 관하여 전반적으로 다루는 분야이다(Russell & Norvig, 2020). AI는 외부 데이터를 올바르게 해석하고 이러한 데이터에서 학습하여, 학습한 내용에 대한 유연한 적응을 통해 특정 목표와 작업을 달성하는 기계 또는 시스템의 능력을 의미한다(Haenlein & Kaplan, 2019). 특히, AI가 모형화(Modeling)하려는 대상은 인간으로서, 기계를 통해 사람의 사고와 행동을 설명하기 위한 시도이며, 이를 위해서는 입력된 프로그램을 자율적으로 확장할 수 있는 기계의 능력이 중요하다(Schank, 1991). 이러한 AI 기술을 통해 기계는 참여하는 작업을 학습함으로써, 시간이 지남에 따라 성능을 스스로 향상할 수 있다(Ayoub & Payne, 2016).

AI 기술은 활용 범위 및 발전 수준에 따라 범용 AI와 특정 영역 중심 AI로 구분할 수 있다. 좁은 의미의 모듈식 AI는 특정 영역에 대한 학습을

통하여 기계 스스로 특정 작업에 관한 성능을 자율적으로 향상할 수 있는 시스템으로서 약한 AI(Weak AI)라고 하지만, 범용 AI는 의미와 가치에 대한 전반적인 이해가 필요한 문제를 포함하여 훨씬 더 광범위하고 추상적인 문제를 해결하기 위해 훨씬 더 유연하게 지식을 활용할 수 있는 시스템으로서 강한 AI(Strong AI)라고도 한다(Ayoub & Payne, 2016). AI 분야는 시스템이 자율적 에이전트 역할을 하고 독립적으로 학습하여 환경을 평가하며 가치, 동기 및 감정 등을 통해 추론하는 능력을 갖춘 새로운 수준의 컴퓨팅으로 이어지고 있다(Barth & Arnold, 1999).

챗봇을 활용하는 과정에서의 문제는 이미 진행되고 있다. 2020년 12월에 스타트업 업체인 ‘스캐터랩(Scatterlab)’이 AI 캐릭터 챗봇으로 ‘이루다(Leeluda)’를 출시하자 특정 이용자 계층을 중심으로 이 캐릭터에게 성차별적 학습을 시킴으로써 사회적 문제가 되었다(연합뉴스, 2021).

AI 기술은 거짓되거나 차별적인 정보를 유통하는 데에만 활용되는 것이 아니라 후원금 모금, 타게팅 선거운동, 유권자 분석 등 여론의 동향을 파악하는 데에도 높은 활용도가 있다. 최근 개인 추천화 알고리즘이 포함된 SNS, 유튜브, 메타 등과 함께 ChatGPT, 바드, 빙 등 생성형 AI에 대한 개인정보 수집 현황을 조사한 결과, 개인 프라이버시와 같은 민감한 정보가 대량 수집되고 있는 것으로 확인되었다. 개인의 검색 기록, 방문 기록, 소셜 미디어 활동 등을 분석하여 개인의 선호도와 관심사를 파악할 수 있으며, 이를 통해 개인의 정치 성향도 추론할 수 있다(세계일보, 2023나). 이처럼 AI 기술은 다양한 기술적 문제를 해결하여 윤리적 및 사회적 문제에 대해 올바르게 추론할 수 있는 능력의 향상이 여전히 필요하다(de Sousa, de Melo, Bermejo, Farias, & Gomes, 2019).

4) AI 리터러시 격차

생성형 AI에 대한 인식도 조사 관련 두 가지 설문조사를 비교하고자 한다. 사회적가치연구원의 2023년 5월 설문조사 결과, 사용한 경험이 있는 비율이 31.1%였다(사회적가치연구원, 2023). 한편, 2023년 8월 리서

치앤리서치의 설문조사를 의뢰한 이강호(2023)에 의하면 생성형 AI를 '사용한 경험이 있다'라고 응답한 비율은 51.5%였으며 '사용한 경험이 없다'라고 응답한 비율은 27.6%, '무엇인지 모른다'라고 응답한 비율은 21.3%다(이강호, 2023). 생성형 AI의 사용 실태는 연령, 학력, 직업, 소득 수준 등에 따라 다양한 양상을 보인다. 대체로 연령이 낮은 그룹에서 생성형 AI 유경험자가 더 많이 나타나는 경향을 보였으며, 생성형 AI의 사용 여부와 학력 및 소득 수준 간에도 밀접한 관련성이 있다. 학력별로는 대학원 이상의 고학력 집단에서 생성형 AI 유경험 비율이 가장 높게 나타난다. 해당 집단의 경우 약 72.5%가 생성형 AI를 사용해 본 적이 있다고 응답하였으며, 그다음으로는 대학교 졸업자 그룹에서 약 54.2%의 유경험률을 기록하고 있다. 반면, 전문대학 졸업자와 고등학교 졸업 이하의 저학력 그룹에서는 상대적으로 낮은 유경험 비율을 보였는데, 각각 약 36.9%와 35.9%의 응답자가 생성형 AI를 사용해 본 적이 있다고 밝히고 있다. 소득이 높을수록 생성형 AI를 경험한 비율도 높아지는 경향을 보인다. 특이한 부분으로 AI의 신뢰도와 관련해서는 연령별로 두드러진 인식 차이를 보인다. 18~29세와 30대는 AI가 정확하다는 응답 비율이 각각 45.4%, 42.2%로 전체 평균보다 낮다. 반면 30대와 40대는 AI가 정확하다고 응답한 비율이 각각 53.9%, 55.7%로 전체 평균보다 높다(이강호, 2023 : 머니투데이, 2023).

이강호(2023)에 의하면 AI의 환각 현상(Hallucination)을 이해하거나 경험하였을 때 AI 답변의 정확성을 낮게 응답한 것으로 분석되었으나 앞으로 생성형 AI 모델 간 경쟁 심화로 기술 발전이 가속화되고 있고 환각 현상이 줄어들고 있으므로 향후 정확성 인식도는 높아질 것으로 예상하였다(머니투데이, 2023).

사회적가치연구원(2023) 설문조사 결과에 의하면 수도권 거주자들의 AI 리터러시는 상대적으로 높은 편이지만 지역간 격차가 존재하는 것으로 파악되었고 AI 리터러시의 격차는 앞으로 더 커질 가능성이 있는 것으로 분석된다. AI를 사용해 보았고, 잘 알고 있다고 응답한 경우(중립 의견 제외)는 AI 리터러시 수준이 높은 사람으로 파악되며 이들 중 AI를 통해 업

무 능력이 향상될 것으로 기대하는 비중은 97.4%로 나타났다. 또한 98.1% 비율로 응답자 대다수는 미래에도 계속해서 AI를 사용할 의사가 있다고 밝혔다(사회적가치연구원, 2023). 반면, AI를 사용한 경험이 없고, 잘 모른다고 응답한 경우는 AI 리터러시 수준이 낮은 사람으로 파악되며 이들 중 65.9%는 AI가 업무에 도움이 될 것으로 생각했지만, 앞으로도 사용할 의향은 50%에 불과했다. 이는 현재 AI 기술에 대한 이해도와 경험에 따라 향후 활용 가능성이 확연히 달라질 것으로 분석된다(사회적가치연구원, 2023). 두 조사에서 공통으로 확인된 사실은 남성과 연령이 낮을수록, 그리고 학력이 높을수록 AI를 활용한 경험이 많고 AI 리터러시 수준이 높다는 것이다.

AI 리터러시 수준을 높고 낮음의 상대적인 부분을 강조하기 위해 사회적가치연구원의 설문 결과를 보다 상세하게 분석하였다. AI 리터러시 수준이 높은 사람들은 전체 1,000명 중 194명으로 AI 기술에 대한 이해도와 사용 경험이 풍부한 사람들로 조사되었고 중립 의견은 제외하였다. AI 리터러시 수준이 낮은 사람은 AI 기술을 잘 알지 못하고, 사용 경험이 없는 경우로 83명으로 조사되었다.

AI 리터러시 수준을 분석한 결과, 남성이고 40대 이상에 고학력이며 수도권에 거주하는 사람일수록 AI 경험이 많고 유용하다고 평가하였고 앞으로도 활용할 가능성도 높게 나타났다. 반면, 활용 경험이 적고 불편함을 느낀 사람들은 앞으로도 이용할 가능성이 크게 떨어지는 것으로 나타났다. 사회 계층별 AI 리터러시 격차는 지속적으로 커질 것으로 전망된다.

따라서, AI 기술이 널리 보급되면서 AI의 순기능과 역기능에 대한 시민들의 인식과 AI 리터러시 수준이 사용 의도에 미치는 영향에 대한 실증적 연구가 필요한 시점이다. AI 기술이 다양하게 사용될 지능정보화시대를 맞이하여 기술의 연속성과 문화의 지각된 인식 수용성 증대 측면에서 다양한 정책을 제시하고 있다(성욱준, 황성수, 2017). 특히 이용자의 수용성 제고 측면에서 모범사례를 창출하여 순기능을 확산하고 프라이버시와 AI의 책임 소재의 문제와 같은 역기능 개선의 필요성을 강조하였다(성욱준, 황성수, 2017). 그러나 AI 기술의 도입에 관한 선행연구에서는 공공

정책에 AI 기술 도입의 필요성을 강조하는 논의에 그쳐 개인의 인식과 기술적 역량이 미치는 영향에 관한 실증적 연구가 부족하다(Sharma, Yadav, & Chopra, 2020; de Sousa et al., 2019).

ChatGPT와 Chatbot을 포함한 AI 기술은 이미 우리 생활의 여러 측면에 널리 퍼져 있다. 이러한 기술은 개인화, 효율성, 가용성, 확장성 등 다양한 편의를 제공하면서 사람들과의 상호작용을 촉진한다. 그러나 현재 기술의 한계 및 자연어처리 기술 부족으로, 정보의 정확도를 높이기 위해 개선해야 할 문제점들이 아직 남아있으며 특히 제한된 언어 능력, 공감 부족, 유지 관리의 필요 그리고 프라이버시 및 사이버 보안 문제 등 다양한 윤리적·사회적 문제점이 존재한다(이철승, 백혜진, 2023). 최근 ChatGPT의 인기가 높아지면서 다양한 분야에서 이 기술을 활용하려는 움직임이 나타나고 있다. 미래 사회에 대비하기 위해서는 부정적인 문제들을 올바르게 판단하기 위해 개인의 AI 리터러시 역량이 더욱 필요하며, AI 리터러시 수준 향상을 통한 ChatGPT 활용은 이용자의 만족감을 더욱 높일 수 있을 것이다.

AI는 나노 기술, 생체공학, 3D 프린팅 등 기술혁명을 가능하게 해주는 여러 요소 가운데 현재 시점에서는 AI 기술이 가장 두드러진다. AI 기술은 가장 큰 편익을 가져다줄 수 있지만, 동시에 개개인의 인간 및 전체 사회에 가장 위협적인 존재가 될 수도 있다. 왜냐하면 AI는 인간 영역의 깊은 곳까지 침투하여 인간이라는 것이 무엇인지, 그리고 우리가 장차 어떤 존재가 될 것인지 등에 대해 근본적인 여러 질문을 제기하기 때문이다. AI 기술은 이제 현실에 더욱 가까워졌다. 최근 몇 달 동안이 가장 많이 발전했으며, 처음에는 일시적인 유행으로 여겨졌지만, 이제는 각 분야에 상당한 영향을 미치고 있다. 최근 통신안내원이 AI로 대체되는 등 사회문제를 발생시켰고, 부정적인 뉴스가 많아졌다.

본 연구에서는 개인의 AI 리터러시 수준과 AI에 대한 기술수용모델을 바탕으로 AI 기술의 혜택과 우려에 대한 지각이 AI 기술의 사용 의도에 미치는 영향을 실증적으로 살펴본다. 이를 위해 사용용이성과 지각된 유용성을 변인으로 한 기술수용모델, AI 리터러시, 그리고 사용의도에 관한

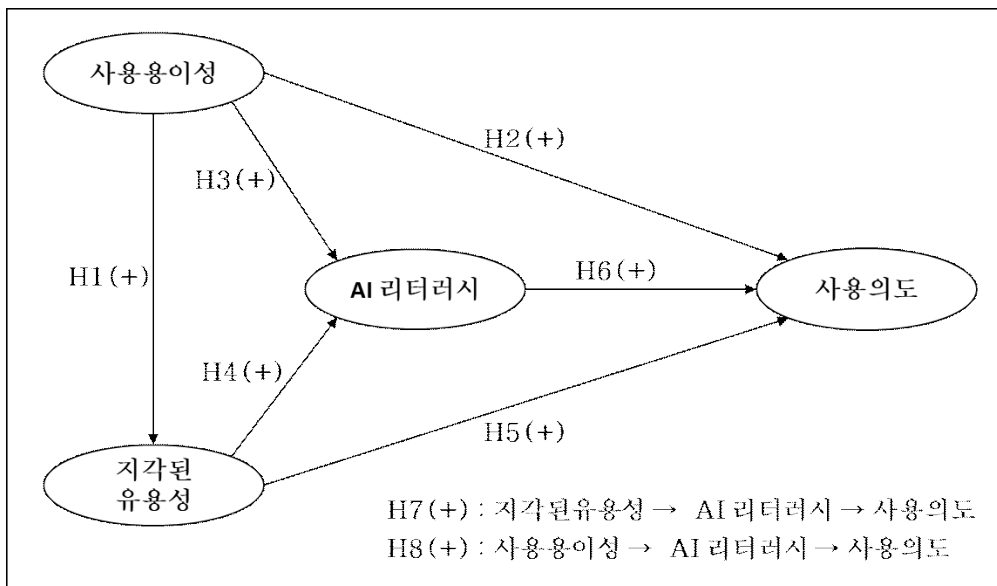
선행연구를 토대로 개념적 분석 모형과 연구가설을 설정한다. 인터넷 설문조사 결과를 이용하여 실증분석을 수행하고자 한다. 실증분석을 통해 본 연구의 결론 부분에서는 분석 결과를 요약하고, 오늘날 사는 평범한 우리들에게 AI 기술이 미칠 수 있는 긍정과 부정의 영향에 대해 연구하여 AI 리터러시 격차를 해소하기 위한 정책 제언의 기초를 제공하고자 한다.

제 3 장 연구설계

제 1 절 연구모형 및 연구가설 설정

1) 연구모형

본 연구는 선행연구를 토대로 AI 기술을 사용하고 있는 사용자에게 관련된 요인과 그 수용의 결정요소를 측정하는 기술수용이론을 바탕으로 하였다. AI 기술이 사용의도에 미치는 영향에 있어서 AI 리터러시 수준의 매개효과를 분석하기 위해 연구가설을 설정하고, 이를 검증하기 위해 [그림 3-1]과 같이 연구모형을 구성하였다. 이후 구조방정식 모형을 활용하여 경로분석을 실시한다. 구조방정식 모형은 측정변수 또는 잠재변수 간의 인과관계에 대해 이론에 기반하여 특정 가설을 평가하기 위한 자료 분석 방법의 하나이다(Mueller & Hancock, 2010). 이를 통해 AI 기술 수용에 영향을 미치는 요인들을 실증적으로 분석한다.



[그림 3-1] 연구모형

기술수용모델은 다양한 과학기술과 제도의 특성에 따라 확장된 기술수용모델로 응용되었다(고영화, 임춘성, 2021). 신기술의 특징에 따라 이용의도에 영향을 미치는 외부 변수 또한 달라질 수 있다(장한진, 노기영, 2017). 이에 따라 본 연구에서는 기술수용모델 이론에 근거하여 사용용이성, 지각된 유용성의 변수를 독립변수로 설정하였다. 신기술인 AI 기술의 수용성을 확인하기 위해 새로운 외부 변수인 AI 리터러시를 매개변수로 활용하여 개념적 연구모형을 설계하였으며, 이후 연구가설을 설정하고 실증적 분석을 통해 검증한다(Ajzen, 1991; Davis, 1989).

2) 연구가설 설정

본 연구의 연구모형에서 설정한 다양한 변수의 관계를 검증하기 위해 AI 기술의 사용용이성, 지각된 유용성, AI 리터러시, 그리고 AI 기술의 사용의도에 대한 연구가설을 설정하였다. 기술이 적용된 서비스의 혜택과 우려에 대해서도 디지털 리터러시 수준에 따라 다르게 인식될 가능성이 있고, 온라인 서비스 활용의 혜택에 관한 인식과 데이터 추적에 따른 프라이버시에 관한 손실의 위험에 관한 인식이 모두 증가할 가능성이 있다(Lutz, 2019). 기술수용모델의 주요 변수인 사용용이성과 지각된 유용성을 독립변수로 설정하고, 리터러시 수준은 태도에 영향을 미치는 경험적 요인 및 학습과도 관련이 있으므로 개인의 AI 기술에 대한 이해도 및 경험에 대한 잠재변수로 AI 리터러시를 설정하여 매개변수로서 사용의도와 의 관계를 검증하기 위해 다음과 같은 연구가설을 설정하고 실증적 분석을 통해 검증한다(Ajzen & Fishbein, 1975; Venkatesh, Morris, & Davis, 2003).

가) AI 기술의 기술수용모델에 관한 가설

기술수용모델의 변수 간의 관계에 있어, AI 기술에 대한 사용용이성이 높아질수록 지각된 유용성이 높아지며, 그 기술의 사용의도가 높아질 것

이라는 정(+)적인 관계를 가정하였다. 또한 AI 기술에 대한 사용용이성과 지각된 용이성이 높아질수록 AI 리터러시 수준도 높아질 것이라는 정(+)
적인 관계를 가정하였다. 그리고 지각된 유용성과 AI 리터러시 수준이 높
아질수록 사용의도가 높아질 것이라는 정(+)
적인 관계를 가정하여 설정
한 연구가설 H1, H2, H3, H4, H5, H6는 아래와 같다.

가설1(H1). AI 기술의 사용용이성은 지각된 유용성에 정(+)
의 영향을 미칠
것이다.

가설2(H2). AI 기술의 사용용이성은 기술 사용자의 사용의도에 정(+)
의 영
향을 미칠 것이다.

가설3(H3). AI 기술의 사용용이성은 AI 리터러시 수준에 정(+)
의 영향을
미칠 것이다.

가설4(H4). AI 기술의 지각된 유용성은 AI 리터러시 수준에 정(+)
의 영향을
미칠 것이다.

가설5(H5). AI 기술의 지각된 유용성은 AI 기술 사용의도에 정(+)
의 영향을
미칠 것이다.

가설6(H6). AI 리터러시 수준은 AI 기술 사용의도에 정(+)
의 영향을
미칠
것이다.

나) AI 리터러시의 매개효과에 관한 가설

다음으로 지각된 유용성과 사용용이성이 높아질수록 사용의도가 향상
되는데 있어서 AI 리터러시 수준이 이를 매개할 것이라는 매개변수의 매
개효과 검증을 위한 연구가설 H7, H8을 아래와 같이 확인할 수 있다.

가설7(H7). AI 리터러시 수준은 AI 기술의 지각된 유용성과 사용의도 간에
매개작용을 할 것이다.

가설8(H8). AI 리터러시 수준은 AI 기술의 사용용이성과 사용의도 간에 매개
작용을 할 것이다.

제 2 절 변수의 조작적 정의

본 연구에서는 기존 기술수용모델을 기반으로 한 선행연구를 참고하여 신뢰도와 타당도가 검증된 바 있는 측정 항목들을 본 연구 방향에 맞게 조정하여 사용하였다. 독립변수인 사용용이성, 지각된 유용성의 기술수용모델의 주요 변수들을 정의하였고 매개변수로는 AI 리터러시를 설정하였다. 그리고 종속변수는 사용의도를 정의하였다. 선행연구를 통해 이러한 변수 간의 관계를 검증 조사하였고 [표 3-1]과 같이 정리하였다.

[표 3-1] 변수의 조작적 정의 및 선행연구

요인	조작적 정의	선행연구
사용용이성	잠재적 이용자가 AI 기술을 이용하는 것이 정신적으로 신체적으로 수고가 적게 들 것이라고 믿는 정도 또는 큰 노력 없이 새로운 기술을 사용할 수 있을 것으로 기대하는 정도	임성진(2012) Rogers(2003), Rai,Lang,&Welker (2002)
지각된 유용성	AI 기술의 효과성에 대한 사용자의 지각된 평가. 잠재 이용자가 AI 기술을 이용하는 것이 직무성과를 향상시킬 것이라고 믿는 정도	윤태환,왕새롬(2023), 김윤경(2022), 이한신,김판수(2019)
AI 리터러시	AI 기술과 각종 응용 프로그램을 이해, 활용하는 데 필요한 지식, 기술 및 윤리적 인식을 포함하는 포괄적 개념을 의미하며, 기술적 이해를 넘어 비판적 사고, 윤리적 고려 및 AI의 광범위한 사회적 영향에 대한 인식을 포함하는 다차원적 개념	김성희(2023), 고윤정(2022), 이철승,백혜진(2022)
사용의도	AI 기술(생성형 AI)을 지속적으로 사용하고자 하는 의도	박유진,이상원(2022), 장창기,성욱준(2022)

제 3 절 설문자료 수집

본 연구에서는 AI 기술에 대한 사용자들의 수용 인식을 파악하기 위하여 전국의 20세 이상 성인을 대상으로 온라인 설문조사를 실시하였다. 설문조사는 온라인 설문조사 전문기관인 (주)더브레인에 의뢰하였으며 2024년 11월 28일부터 12월 6일까지 9일간 단순 무작위 표본 추출 방법을 통해 진행되었다. 그 결과, 총 340명의 데이터가 수집되었고 불성실한 응답 28개를 삭제하여 최종적으로 총 312개의 데이터가 분석에 활용되었다.

특히 2023년부터 우리 사회에서 큰 파급력을 미치고 이슈로 떠오른 AI의 영향력에 대한 설문 내용을 포함하고 있다. 본 설문은 전국 17개 시도에 거주하는 성인 남녀 312명을 대상으로, 구조화된 설문지를 이용한 온라인 응답 조사 방식으로 진행되었다. 설문 조사 개요는 다음 [표 3-2]와 같이 정리하였다.

[표 3-2] 설문 조사 개요

항목	내용
조사기간	2024년 11월 28일 ~ 12월 6일
조사방법	구조화된 설문지를 이용한 온라인 설문조사
조사대상	전국 17개 시도 지역에 거주하는 성인 남녀 312명
표본오차	모집단 340명 중 312명을 표본으로 선정하여 실시되었으며, 표본오차는 $\pm 2.23\%p$ (95% 신뢰수준, 표집방법으로 단순 무작위추출 방식을 활용)

제 4 절 설문지 구성

본 설문지에서 설문 대상자를 위해 설정 AI 기술의 정의, 종류, AI 리터러시에 대한 기본 개념은 다음과 같다.

AI 기술(생성형 AI : Generative AI)이란 사용자의 특정한 요구사항에 따라 결과물을 능동적으로 생성하는 인공지능(AI) 기술을 의미하며, 입력된 데이터를 학습하고, 학습된 데이터를 기반으로 새로운 데이터 결과를 만들어내는 기술을 의미한다.

AI 기술(생성형 AI)을 사용해 본 경험에 대한 범주는 ChatGPT(Open AI), 바드(Bard, 구글), Copilot(마이크로소프트), Claude(Anthropic), Cu e(네이버), 뤼튼, 미드저니(Midjourney), DALL-E, Perplexity 등과 같은 인공지능, 그 외 생성형 AI를 사용한 적이 있는 사람이 포함된다. 이번 표 본 312명 중 3명은 사용 경험이 없는 응답자였다.

AI 리터러시에 대한 설명은 AI 기술과 각종 응용 프로그램을 이해, 활용하는 데 필요한 지식, 기술 및 윤리적 인식을 포함하는 포괄적 개념을 의미하며, 기술적 이해를 넘어 비판적 사고, 윤리적 고려 및 AI의 광범위한 사회적 영향에 대한 인식을 포함하는 다차원적 개념을 의미한다.

본 연구는 최근 사용량이 확산되고 있는 AI에 대해 기술수용모델을 기반한 AI 사용의도에 미치는 영향을 확인하고자 설문 문항을 구성하였다. 설문조사 항목 중 사용용이성, 지각된 유용성을 독립변수로 설정하고 12개 문항, 사용의도를 종속변수로 선정하고 5개 문항, AI 사용의도를 유발하는 데 있어 AI 리터러시 수준이 긍정적인 매개 역할을 하는지도 확인하고자 AI 리터러시를 매개변수로 12개 문항을 설정하였고 리커트 5점 척도로 구성하였다. 설문 응답자의 특성을 파악하기 위해 일반적인 인구통계학적 특성(성별, 연령, 소득, 직업, 학력, 거주지의 6개 문항)과 일반 이용 경험 유무 1개 문항을 함께 활용되었다. 설문지에 제시된 문항은 본 논문의 부록에 첨부하였다.

AI 기술이 도입되어 일반대중에게 본격적으로 소개되는 시점에서, AI에 대한 사용용이성, 지각된 유용성, 사용의도 그리고 이 신기술이 우리

사회에 미칠 변화에 대해 복합적으로 측정한 데이터는 많지 않다. 따라서 AI 리터러시 수준이 기술 사용 의도에 미치는 영향을 실증적으로 분석 검증하는 본 연구가 적합하다고 판단하였다. AI 기술에 대한 인식과 사용 의도를 조사하기 위해 설문 항목을 구성하였고, 사전 조사를 통해 설문 항목의 타당도를 확인한 후 본 조사를 실시하였다.

제 5 절 분석방법

본 연구에서는 수집된 자료 분석을 위해 다음과 같은 통계 분석 방법을 사용하였다. 분석을 위해 SPSS 27.0 프로그램을 사용하였으며, 먼저 연구 대상자의 인구통계학적 분석을 위해 빈도분석을 실시하였으며, 연구 변수들의 기술 통계량(최소값, 최대값, 평균, 표준편차, 첨도, 왜도) 산출을 통해 자료의 기본적인 특성을 파악하였다. 또한 주요 변수들에 대해 피어슨(Pearson) 상관분석을 실시하였다.

다음으로 구조방정식모델을 활용하여 구성 변수 간의 관계를 분석하였다. 확인적 요인분석을 통해 측정항목과 모형의 타당도 평가를 시행하였으며 간접효과를 확인하기 위해 경로분석을 하였다. 이와 같은 구조방정식 분석을 위해서 AMOS 29.0 프로그램을 사용하였다. 모든 통계 분석은 유의수준 .05에서 실시하였으며, 구체적인 p값과 효과 크기 등을 함께 제시하였다. 위와 같은 방식으로 분석 방법을 체계적이고 구체적으로 기술함으로써 연구의 분석 절차와 결과를 이해하기 쉽게 전달하고자 하였다. 자료 통계 분석방법에 대한 요약은 [표 3-3]과 같다.

[표 3-3] 통계 분석 방법

구분	분석 방법	분석 내용
표본의 특성	빈도분석	표본의 인구통계학적인 내용
변수의 특성	기술통계분석	측정 항목별 평균, 표준편차, 왜도, 첨도
	신뢰도 및 타당도 분석	신뢰도(크론바하 알파), 개념 타당도
변수 간 관계	상관관계 분석 (Pearson)	변수 간 상관관계
변수의 타당도와 신뢰도	확인적 요인분석 (CFA)	모형적합도, 집중타당도, 구성개념신뢰도(CR), 평균분산추출(AVE)
가설검증	구조방정식 (CB-SEM)	변수 간 종합적 인과관계 분석 (경로분석, 매개효과분석)

제 4 장 연구결과

제 1 절 인구통계학적 특성 및 기술통계 분석

1) 연구대상자의 인구통계학적 특성

연구 대상자의 인구통계학적 특성은 [표 4-1]과 같이 나타났다. 먼저 설문조사 대상 312명의 성별 분포는 남성이 53.2%, 여성이 46.8% 비율로 조사되었으며 남성이 20명이 더 많았다. 연령대는 20~30대가 152명으로 48.7%, 40대 이상이 160명으로 51.3% 40대 이상이 약간 더 많았다. 본 연구에서는 상대적 비교를 강조하기 위해 학력을 2개 영역으로 SPSS 코딩 변환하였다. 대졸 이상의 고학력자는 233명으로 74.7%, 대졸 미만의 기초학력자는 79명, 25.3%로써 대부분 참가자가 대학 교육을 받은 것으로 나타났다. 거주지역도 상대적 비교를 강조하기 위해 2개 영역으로 SPSS 코딩 변환하여 구분하였으며 참가자의 237명, 76.0%는 수도권 지역(서울, 경기, 인천), 75명, 24%의 참가자는 비수도권에 거주 중인 것으로 분석되었다.

직업의 경우 사무직이 167명(53.5%)으로 가장 많았으며, 자영업 34명(10.9%), 전문직 29명(9.3%), 관리직 26명(8.3%), 판매/서비스업 23명(7.4%), 학생 20명(6.4%) 순으로 나타났다.

가구 월수입의 경우, 700만원~800만원 미만 146명(46.8%), 900만원 이상 107명(34.3%)으로 가장 많았다. 500만원~700만원 미만 33명(10.6%), 500만원 미만 26명(8.3%) 순으로 나타났다.

[표 4-1] 표본의 인구통계학적 특성(n=312)

구분	분류	빈도	비율(%)
성별	남성	166	53.2
	여성	146	46.8
연령	20~39세	152	48.7
	40세 이상	160	51.3
교육수준	기초학력자	79	25.3
	고학력자(대졸 이상)	233	74.7
거주지역	수도권(서울/경기/인천)	237	76.0
	비수도권	75	24.0
직업별	농업/임업/수산업	0	0
	자영업	34	10.9
	판매/서비스업	23	7.4
	기술직	3	1.0
	사무직	167	53.5
	경영 관리직	26	8.3
	전문직	29	9.3
	자유직	2	.6
	전업주부	1	.3
	학생	20	6.4
	무직/기타	7	2.2
월수입 (만원)	~299	10	3.2
	300~499	16	5.1
	500~699	33	10.6
	700~899	146	46.8
	900~	107	34.3
Total		312	100.0

2) 주요 변수의 기술통계량

본 연구에서 설정한 가설과 연구 문제를 검증하기 전 주요 변인들에 대한 기초적인 기술통계량을 확인하였다. 모든 변인은 각 측정항목의 평균을 활용하여 대푯값을 구하여, 기술통계 분석을 실시하였다. 분석 결과 모든 개념의 평균이 3 이상으로 나타났으며, 모든 측정 문항의 평균 및 표준편차는 다음 [표 4-2]와 같이 나타났다.

[표 4-2] 주요 변수의 측정 항목별 평균 및 표준편차

변수	설문문항	평균	표준 편차
사용 용이성	1. AI의 이용 방법은 쉽다고 생각한다	3.50	.95
	2. AI 이용에 금방 익숙해질 수 있다고 생각한다	3.50	.94
	3. AI의 방법은 명확하고 간단하다고 생각한다	3.48	.89
	4. AI는 내가 원하는 것을 쉽게 할 수 있게끔 해준다	3.47	.94
	5. AI를 사용하면 문제해결을 쉽게할 수 있다	3.50	.91
지각된 유용성	1. AI를 이용하면 내 목적을 더 빠르게 달성할 수 있다고 생각한다	3.43	1.11
	2. AI를 이용하면 내 관심사나 업무를 효율적으로 처리하는데 도움이 된다	3.47	1.12
	3. AI를 사용하면 유용한 정보를 얻을 수 있다고 생각한다	3.36	1.16
	4. AI를 사용하면 생산성이 증가한다	3.52	1.12
	5. AI는 전반적으로 나에게 유용하다고 생각한다	3.47	1.14
AI 리터 러시	1. 나는 AI가 응답한 정보를 이해할 수 있다	3.58	1.04
	2. 나는 AI가 무엇인지 설명할 수 있다	3.54	1.06
	3. 나는 어떻게 AI를 이용해 내게 도움이 되는 정보를 요구할 수 있는지 알고 있다	3.48	1.10
	4. 나는 AI가 응답한 정보가 믿을만 한지 않은지 판단할 수 있다	3.40	1.09
	5. 나는 AI를 통해 얻은 정보가 나의 상황에 적용가능한지 선별할 수 있다	3.47	1.11
	6. 나는 AI를 통해 목적에 맞는 정보를 수집하고 판단할 수 있다	3.44	1.09
	7. 나는 AI를 통해 얻은 정보를 필요한 순간에 적절하게 활용할 수 있다	3.54	1.08

	8. 나는 AI를 통해 나에게 가장 도움이 되는 정보가 무엇인지 구분할 수 있다	3.40	1.12
	9. 나는 AI를 통해 빠른 시간 내에 필요한 정보를 찾을 수 있다	3.53	1.08
	10. 나는 AI를 통해 다양한 생각과 의견을 검색할 수 있다	3.49	1.09
	11. 나는 AI 결과물에 대해 그대로 받아들이는 것이 아니라 해석이 필요하다는 것을 이해하고 있다	3.57	1.14
	12. 나는 AI가 응답한 정보가 거짓이거나 오류일 수 있다는 것을 이해하고 있다	3.59	1.06
사용의도	1. AI를 이용하면 내 목적을 더 빠르게 달성할 수 있다고 생각한다	3.67	1.65
	2. AI를 이용하면 내 관심사나 업무를 효율적으로 처리하는데 도움이 된다	3.67	1.63
	3. AI를 사용하면 유용한 정보를 얻을 수 있다고 생각한다	3.71	1.64
	4. AI를 사용하면 생산성이 증가한다	3.66	1.65
	5. AI는 전반적으로 나에게 유용하다고 생각한다	3.64	1.64

본 연구에서 AI의 기술수용모델의 핵심 변수와 매개변수에 대한 정규성을 분석하기 위해 전체 변수에 대한 기술통계 분석을 실시하였으며, 그 결과는 [표 4-3]에 제시한 바와 같다.

[표 4-3] 측정변수의 기술통계 분석

변수	최소값	최대값	평균	표준편차	왜도	첨도
사용용이성	1.2	5.0	3.49	.79	-.147	-.154
지각된 유용성	1.0	5.0	3.45	.98	-.385	-.297
AI 리터러시	1.0	5.0	3.50	.91	-.501	-.193
사용의도	1.0	5.0	3.67	1.62	-.823	-1.029

변수별로 보다 상세한 분석은 다음과 같이 나타났다. 첫째, 사용용이성에 대한 평균은 3.49(SD=.79)로 나타났다. 이를 5점 리커트 척도의 평균인 3.5와 비교해 보면 전체 평균 수준인 것으로 확인되었다. 둘째, 지각된 유용성에 대한 평균은 3.45(SD=.98)로 평균의 수치로 나타났다. 셋째, AI 리터러시에

대한 설문 문항에서 참가자의 평균은 3.50(SD=.91)로 전체 평균과 동일 수준으로 조사되었다. 넷째, 사용의도에 대한 참가자의 평균은 3.67(SD=1.62)로 네 가지 변수 중 평균 대비 가장 높은 수준으로 확인되었다. 이는 AI가 우리 생활 전반적으로 점차 스며들고 있으며, 많은 사람이 이를 적극적으로 활용하고자 하는 의지가 높다는 것을 보여준다.

변수들의 정규성을 검증 확인하기 위해 왜도와 첨도를 [표 4-3]과 같이 분석하였다. 왜도의 절댓값이 3.0 이하이고, 첨도의 절댓값이 8.0 이하일 때에는 해당 변수가 정규분포를 따른다고 볼 수 있다. 변수들의 왜도와 첨도를 분석한 결과, 모든 변수의 왜도와 첨도는 절댓값 2 미만으로 정규분포를 따르는 것으로 확인되었다.

제 2 절 신뢰도 및 타당도 분석

신뢰도 및 타당도 분석결과는 [표 4-4]와 같다. 분석결과, 각 설문문항 요인 항목들의 요인 적재치는 모두 0.5 이상으로 개념 타당도가 확보되었으며, 신뢰도는 크론바하 알파 계수(Cronbach's Alpha) 값이 모두 0.6 이상으로 요인 항목 간의 내적 일관성이 확보되었다.

본 연구에서는 설문문항의 타당도를 분석하기 위해 요인분석을 실시하였고 KMO(Kaiser-Meyer-Olkin)와 Bartlett(Bartlett's Test of Sphericity) 방식을 사용하였다. KMO는 변수 간의 편상관을 확인하는 것으로 변수의 숫자와 케이스의 숫자의 적절성을 나타내는 표본 적합도를 의미한다(노경섭, 2019). 일반적으로 $KMO > .5$, 이면 요인분석을 실시하는 것이 적절하다고 판단된다. Bartlett은 요인분석을 할 때 사용되는 상관계수의 행렬이 대각행렬이면 요인분석을 하는 것이 부적절하며 Bartlett 값에서 $p < .05$ 이면 대각행렬이 아님을 의미하므로 요인분석을 하는 것이 적절하다(노경섭, 2019). 본 연구의 KMO 값은 .971로 .5보다 큰 것으로 확인되었고, Bartlett 값에서 $p < 0.01$ 로 확인되어 타당도가 확보된 것으로 확인되었다.

[표 4-4] 신뢰도 및 타당도 분석결과

설문문항	성분				Cronbach's Alpha
	ITU	AIL	PEOU	PU	
ITU3	.982	.706	.618	.668	.992
ITU1	.982	.706	.604	.645	
ITU4	.982	.711	.586	.654	
ITU2	.980	.716	.604	.674	
ITU5	.977	.702	.603	.675	
AIL8	.621	.850	.478	.461	.960
AIL5	.641	.848	.431	.514	
AIL7	.589	.832	.420	.508	
AIL6	.591	.832	.426	.470	
AIL9	.576	.816	.451	.510	
AIL11	.624	.816	.472	.375	
AIL2	.592	.809	.386	.474	
AIL4	.580	.806	.356	.482	
AIL3	.606	.803	.404	.531	
AIL10	.595	.800	.405	.450	
AIL1	.625	.790	.481	.450	
AIL12	.561	.790	.500	.386	
PEOU5	.493	.415	.819	.373	.905
PEOU2	.542	.452	.815	.443	
PEOU1	.514	.408	.813	.456	
PEOU4	.568	.492	.800	.443	
PEOU3	.514	.475	.792	.363	
PU3	.634	.582	.399	.859	.918
PU1	.654	.536	.465	.845	
PU2	.632	.581	.490	.832	
PU5	.657	.509	.581	.773	
PU4	.647	.547	.607	.760	
KMO(Kaiser-Meyer-Olkin)					.971
Bartlett 구형성 검증 (Bartlett's Test of Sphericity)			Chi-Square		9575.132
			df(p)		351(.000)

* ITU(사용의도), AIL(AI 리터러시), PEOU(사용용이성), PU(지각된유용성)

제 3 절 주요 변수의 상관관계 분석

본 연구에서는 연구모형으로 설정된 주요 변수의 상관관계를 검증하기 위해 피어슨의 상관관계 분석(Pearson's correlation analysis)을 실시하였다. 분석한 결과는 [표 4-5]와 같다.

[표 4-5] 상관관계 분석

구분	사용용이성	지각된 유용성	AI 리터러시	사용의도	Cronbach's Alpha
사용용이성	1.000				.905
지각된 유용성	.650**	1.000			.918
AI 리터러시	.577**	.689**	1.000		.960
사용의도	.660**	.788**	.745**	1.000	.992

* $p < .05$, ** $p < .01$, *** $p < .001$

주요 변수들은 모두 정(+)적 상관관계를 보였으며, 변수 간에는 통계적으로 유의미한 상관관계를 나타냈다. 구체적으로 살펴보면, 사용용이성과 지각된 유용성은 $r = .650(p < .01)$ 로 나타났으며, 사용용이성과 AI 리터러시는 $r = .577(p < .01)$ 로 상관관계가 유의미한 것으로 나타났다.

지각된 유용성과 AI 리터러시는 $r = .689(p < .01)$ 로 나타났고, 사용용이성과 사용의도 간에는 $r = .660(p < .01)$, 지각된 유용성과 사용의도 간에는 $r = .788(p < .01)$, AI 리터러시와 사용의도 간에는 $r = .745(p < .01)$ 로 나타났으며 모두 정(+)의 영향을 미치는 것으로 확인되었다. 변수들 간의 상관관계수가 0.8 이상이면 일반적으로 다중공선성의 위험이 크다고 판단되지만, 이번 연구에서는 다중공선성이 의심되는 변수는 발견되지 않았다. 실제 다중공선성 여부를 판단하기 위해 다중회귀분석을 통해 VIF가 10을 넘는지를 확인한 결과 모든 측정변수가 문제가 되지 않는 것이 확인되었다(권순동, 2015).

제 4 절 확인적 요인분석

구조방정식 모형의 적합도를 평가하기 위해선 절대적합지수와 중분적합지수, 간명적합지수를 통해 종합적으로 살펴보아야 한다. 일반적으로 모델적합도 평가에 많이 사용되는 절대적합지수는 χ^2 (카이제곱 통계량), Normed χ^2 (χ^2/df)이 있으며, χ^2 통계량 같은 경우 p값이 0.05 이상, Normed χ^2 의 값이 3 이하여야 모델이 적합하다고 판단하지만, χ^2 값은 표본 크기, 모델 복잡성 등 여러 요인으로 인해 달라질 수 있으므로 χ^2 통계량만을 기준으로 판단하는 것은 적절하지 않다. 이에 따라 다른 지수를 함께 고려해 모델적합도를 판단해야 할 필요가 있다(우종필, 2012; 신건권, 2013). 앞선 통계량과 더불어, 모델적합도를 확인하는 방법으로는 RMR(0.05 이하 우수), RMSEA(0.05 이하 우수, 0.08 이하 양호, 0.1 이하 보통) 등이 있으며, 중분적합지수는 CFI(0.90 이상 우수), TLI(0.90 이상 우수), IFI(0.90 이상 우수) 등이 있다. 간명적합지수로는 AGFI(0.90 이상 우수, 0.80 이상 적합/양호) 등의 값을 확인하는 방법이 있다.

확인적 요인분석 결과 본 연구에서 이용한 측정 모형의 적합지수는 다수의 조건에서 대체적으로 만족할 만한 수준인 것으로 나타났다($\chi^2=514.678$, $df=318$, $p<.001$, Normed $\chi^2=1.618$, RMR=0.038, RMSEA=0.045, IFI=0.979, CFI=0.979, TLI=0.977, AGFI=0.868). 자세한 분석 결과는 [표 4-6]과 같다.

[표 4-6] 확인적 요인분석 모형적합도 분석

모형적합도	적합 기준값	결과	판단
Normed χ^2	3 이하	1.618	양호
RMR	.08이하 양호 / .05이하 우수	.038	우수
SRMR	.08이하 양호 / .05이하 우수	.032	우수
GFI	.8이상 양호 / .9이상 우수	.889	양호
AGFI	.8이상 양호 / .9이상 우수	.868	양호
NFI	.9 이상	.948	적합
TLI	.9 이상	.977	적합
CFI	.9 이상	.979	적합
RMSEA	.08 이하 양호 / .05 이하 우수	.045	우수

경로분석을 실시하기 전 설문문항을 통해 측정하고자 하는 구성개념의 타당도와 신뢰도를 확인하기 위해 확인적 요인분석(CFA: Confirmatory Factor Analysis)을 수행하였다. 확인적 요인분석은 연구자가 이론을 바탕으로 가정한 구성개념과 측정 문항 간의 관계를 지정하고, 그 인과관계를 확인하는 기법이다(Anderson & Gerbing, 1988).

확인적 요인분석은 관측변수와 잠재변수 간의 요인부하량을 측정함으로써, 모델의 전반적인 적합도를 평가할 수 있으므로 구성개념 타당도(Construct Validity)를 측정하는 데 효과적으로 사용될 수 있다(Anderson & Gerbing, 1988). 구성개념 타당도는 특정 구성개념과 이를 측정하는 관측변수 간의 일치도를 나타낸다. 구성개념이 관측변수에 의해 얼마나 잘 측정되었는지를 나타내며, 측정 도구가 측정하고자 하는 구성개념의 값을 정확히 측정하는 정도에 관한 것을 말한다. 판별타당도(Discriminant Validity), 집중타당도(Convergent Validity), 법칙타당도(Nomological Validity)

y) 등으로 분류된다. 확인적 요인분석에서의 추정법은 보편적으로 많이 사용되고 있는 최대우도법(Maximum Likelihood: ML)을 통해 이루어졌다.

집중타당도는 잠재변수를 측정하는 관측변수들이 얼마나 일치하는지를 나타낸다. 이는 측정 항목들이 구성개념을 일관성 있게 측정했을 때 높은 상관관계를 보임으로써 확인된다. 확인적 요인분석을 통해 측정 문항과 구성개념 간의 타당도와 신뢰도를 확인할 수 있다. 즉, 구성개념신뢰도(CR: Composite Reliability)를 통해 구성개념에 대한 측정 문항의 신뢰도를 검증하고, 일반적으로 그 값이 0.7 이상이면 구성개념의 신뢰도를 만족한다(Segars, 1997). 또한 집중타당도와 판별타당도는 평균표본추출(AVE: Average Variance Extracted)과 구성개념 간의 상관관계를 통해서 검증하며, 평균 표본추출(AVE)이 일반적으로 0.5 이상이고 구성개념 간의 상관관계 값의 제곱보다 커야 타당도를 충족한다 (Anderson & Gerbing, 1988; Bagozzi & Yi, 1988; Fornell & Larcker, 1981; Segars, 1997). [표 4-7]에서 나타난 것처럼, 확인적 요인분석에 대한 구성개념 신뢰도(CR)가 모든 구성개념에 대해 기준치인 0.7보다 높은 것(0.898~0.979)으로 확인되고 이는 개념 신뢰도가 확보되었다 볼 수 있다, 평균분산추출도(AVE)의 값이 기준치 0.5보다 높은 것(0.629~0.905)으로 확인되어 집중타당도가 확인되었다. 설문 항목의 변수별 집중타당도 분석 결과는 [표 4-7]과 같다.

[표 4-7] 확인적 요인분석 결과 : 집중타당도

변수	최소값	B	β	S.E	C.R	P	AVE	CR
사용 용이성	사용용이성1	1	.802				.690	.918
	사용용이성2	1.005	.821	.062	16.129	***		
	사용용이성3	.924	.791	.06	15.370	***		
	사용용이성4	1.018	.826	.063	16.273	***		
	사용용이성5	.968	.809	.061	15.818	***		
지각된 유용성	지각된유용성1	1	.820				.638	.898
	지각된유용성2	.971	.810	.058	16.638	***		
	지각된유용성3	1.040	.838	.059	17.503	***		
	지각된유용성4	1.015	.847	.057	17.779	***		
	지각된유용성5	1.000	.843	.057	17.642	***		
AI 리터 러시	AI리터러시1	1	.795				.629	.953
	AI리터러시2	1.040	.808	.064	16.336	***		
	AI리터러시3	1.069	.806	.066	16.278	***		
	AI리터러시4	1.062	.803	.066	16.205	***		
	AI리터러시5	1.144	.851	.065	17.554	***		
	AI리터러시6	1.097	.830	.065	16.959	***		
	AI리터러시7	1.085	.831	.064	16.983	***		
	AI리터러시8	1.153	.850	.066	17.531	***		
	AI리터러시9	1.064	.817	.064	16.592	***		
	AI리터러시10	1.060	.801	.066	16.144	***		
	AI리터러시11	1.124	.814	.068	16.502	***		
	AI리터러시12	1.010	.788	.064	15.789	***		
사용 의도	사용의도1	1	.982				.905	.979
	사용의도2	.985	.981	.015	63.723	***		
	사용의도3	.995	.983	.015	65.226	***		
	사용의도4	.999	.981	.016	63.919	***		
	사용의도5	.989	.978	.016	60.746	***		

[표 4-8]의 대각선 계수는 변수 AVE의 제곱근 값을 나타내고 있으며 해당 변수 AVE의 제곱근 값이 해당 잠재변수 간의 상관계수보다 높으면 판별타당도가 확보된다고 할 수 있다. 즉, 변수 AVE의 제곱근 값이 가장 큰 상관계수보다 크면, 그 구성요소와는 잘 구별되는 것으로 볼 수 있다. 따라서 본 연구의 [표 4-8]에서 검증한 AVE 제곱근 값 또한 판별타당도를 충족하는 조건이다. 분석결과, AVE의 제곱근 값이 모든 다른 변수와의 상관관계 값보다 높게 나타났으며, 이는 측정의 판별타당도와 신뢰도를 확보되었음을 확인할 수 있다(Anderson & Gerbing, 1988).

[표 4-8] 구성개념신뢰도(CR)와 평균분산추출(AVE)

구분	AVE	CR	사용 용이성	지각된 유용성	AI 리터러시	사용의 도
사용용이성	.690	.918	.831			
지각된 유용성	.638	.898	.650**	.799		
AI 리터러시	.629	.953	.577**	.689**	.793	
사용의도	.905	.979	.660**	.788**	.745**	.951

* $p < .05$, ** $p < .01$, *** $p < .001$

a) 음영표기 셀 : AVE의 제곱근 값을 표시함

b) 하단부 셀 : 변수 간 상관계수를 표시함

제 5 절 연구가설 검증

1) 연구모형의 적합도 검증

연구모형의 적합도 검증과 연구가설의 검증을 위해 AMOS 29.0을 사용하여 구조방정식모형 분석을 수행하였다. 모수의 추정법으로는 최대우도법(ML: Maximum Likelihood)을 사용하였다.

사용용이성, 지각된 유용성과 사용의도 간에 AI 리터러시 수준이 간접 효과가 있는지를 살펴보기 위해 경로분석을 실시하였다. 분석결과 경로분석 모형의 적합도는 χ^2 (Chi-square)=514.678(p<.001), df=318, CMIN/df=1.618, RMR=.038, SRMR=.032, GFI=.889, AGFI=.868, CFI=.979, NFI=.948, TLI=.977, RMSEA=.045로 나타났다. 자세한 분석 결과는 [표 4-9]와 같다.

[표 4-9] 연구모형 모형적합도 분석

모형적합도	적합 기준값	결과	판정
Normed χ^2	3 이하	1.618	양호
RMR	.08이하 양호 / .05이하 우수	.038	우수
SRMR	.08이하 양호 / .05이하 우수	.032	우수
GFI	.8이상 양호 / .9이상 우수	.889	양호
AGFI	.8이상 양호 / .9이상 우수	.868	양호
NFI	.9 이상	.948	적합
TLI	.9 이상	.977	적합
CFI	.9 이상	.979	적합
RMSEA	.08 이하 양호 / .05 이하 우수	.045	우수

*p<.05, **p<.01, ***p<.001

연구모형의 모형적합도를 확인하기 위해 χ^2 검증을 할 경우, χ^2 검증은 표본이 크기에 매우 민감하여 표본의 크기가 클수록($n > 200$ 인 경우) 연구 모형은 기각되기 쉬우며, 표본이 $n > 400$ 인 경우 거의 대부분 통계적으로 유의한 것으로 나타난다(히든그레이스, 2023). 통계적으로 χ^2 (Chi-square) p값이 .05보다 작으면, 모델이 데이터에 적합하지 않다고 할 수 있다. χ^2 값은 표본 크기, 모델 복잡성 등 여러 요인으로 인해 달라질 수 있으므로 χ^2 통계량만을 기준으로 판단하는 것은 적절하지 않다. 이에 따라 다른 지수를 함께 고려해 모델적합도를 판단해야 할 필요가 있다(김계수, 2010; 우종필, 2012).

앞선 χ^2 값 통계량과 더불어, 모델적합도를 확인하는 방법으로는 RMR/SRMR(.05 이하 우수), RMSEA(.05 이하 우수, .08 이하 양호, .10 이하 보통) 등이 있으며, 중분적합지수는 CFI(.90 이상 우수), TLI(.90 이상 우수) 등이 있다. 간명적합지수로는 AGFI(.90 이상 우수, .80 이상 적합/양호) 등의 값을 확인하는 방법이 있다. 따라서 본 연구에서는 χ^2 검증에서 표본 크기에 의한 영향력을 최소화하기 위한 대안으로 RMR/SRMR, RMSEA, CFI, TLI, AGFI 모형적합도 지수를 사용하였다(김계수, 2010; 우종필, 2012; 히든그레이스, 2023).

RMR(Root Mean-square Residual)은 일반적으로, .05 이하면 좋은 것으로 설명되고 있으나 연구에 따라 조사자 스스로 수용 가능 수준을 결정하는 것이 적절하며, 대체로 .05~1.0 사이의 값이면 충분하다(김계수, 2010). 본 연구모형의 RMR값은 .038로 이에 부합된다고 할 수 있다.

SRMR(Standardized Root Mean Residual)은 RMR을 표준화한 수치로 .08 미만이 값은 일반적으로 허용되는 기준으로 판단되나 학자에 따라 1.0 이하의 경우 적절하다고 판단하기도 한다. 본 연구의 SRMR의 값은 .032로 수용할 수 있는 수준이다.

GFI(Goodness of Fit Index)와 AGFI(Adjusted Goodness of Fit Index)는 일반적으로 .90 이상이면 우수한 것으로 판단되지만, Netemeyer, Boles, McKee, & McMurrian(1997)에 의하면 GFI=.85, AGFI=.75 이상이 최소한의 적합 기준(Marginal Fit)으로 제안되었는데, 이는 표본의 특

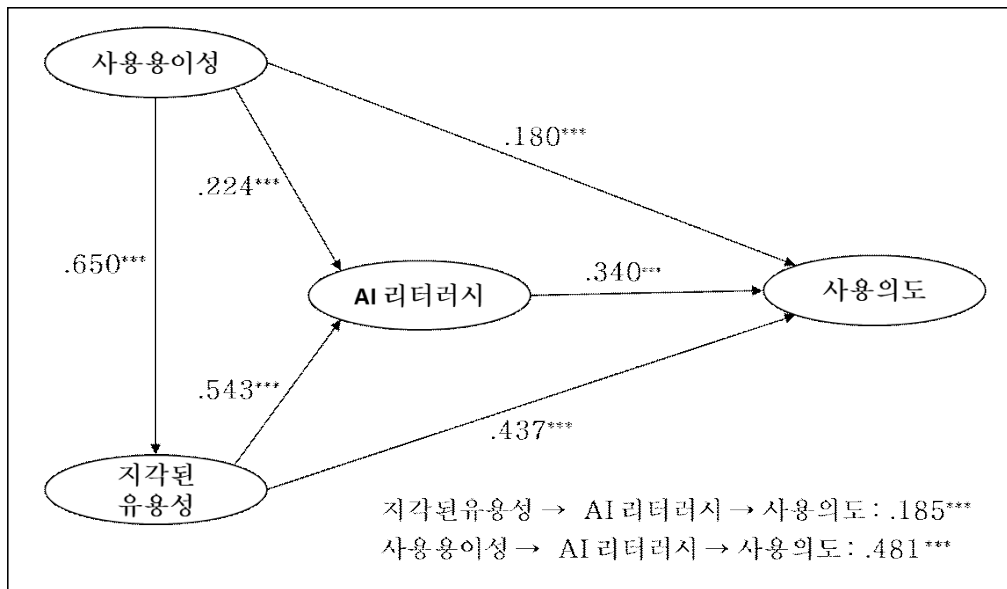
성 때문에 GFI와 AGFI가 낮아질 수 있음을 고려한 것이다. 최근에는 특히 복잡한 모델이나 큰 표본 크기를 가진 모델에서 이러한 높은 기준치를 충족하기 힘들기 때문에 이러한 엄격한 기준은 너무 경직되었다는 지적이 있다. 따라서 본 연구모형의 $GFI=.889$, $AGFI=.868$ 의 경우 최소 기준을 만족하게 하므로 수용 가능한 수준이라고 평가할 수 있다(Netemeyer, Bolles, McKee, & McMurrian, 1997).

NFI(Normed Fit Index)는 증분적합지수의 기본이 되는 적합도로서 .90 이상이면 양호한 것으로 간주한다. 본 연구의 NFI는 .948으로 나타났으며, 이는 영모델(Null Model)에서 연구모형이 94.8% 향상된 것을 의미한다. TLI(Turker-Lewis Index)는 일반적으로 .90 이상이면 양호한 것으로 간주한다. 본 연구의 TLI는 .977로 양호한 것으로 확인된다. CFI(Comparative Fit Index)는 일반적으로 0~1 사이의 값을 가지며, .90 이상이면 좋은 적합도를 가지는 것으로 본다. 본 연구의 CFI는 .979로 적합한 것으로 나타났다. RMSEA(Root-Mean-Square Error of Approximation)는 일반적으로 .10 이하일 경우 모델의 적합도는 보통이고, .05 이하일 경우에 매우 좋으며, .08 이하면 양호한 것으로 판단한다(우종필, 2012). 본 연구의 RMSEA의 값은 .045로 양호한 것으로 판단된다.

이상의 결과들을 종합해 보면, 본 연구모형은 전반적으로 수용할 수 있는 적합도를 보여주었다. 따라서, 이 연구에서는 연구 단위 간의 인과관계를 설명하는 데 큰 문제가 없을 것으로 판단된다.

2) 가설 검증

본 논문에서 제안된 연구모형이 타당하다고 판단되므로, 구조방정식 모델을 통해 설정된 가설들에 대한 경로 분석결과를 살펴볼 수 있다. 각 경로계수의 추정치는 최대우도법(Maximum Likelihood)을 사용하여 추정되었다. 확인적 요인분석 결과를 토대로, 독립변수인 지각된 유용성과 사용용이성과 매개변수인 AI 리터러시가 사용의도에 미치는 영향을 검증하기 위해 경로분석을 실시했다. 분석은 전체 집단에 대한 경로분석을 수행하여 독립 및 매개변수의 직접효과를 확인하고 독립변수에 대한 매개변수의 간접효과를 확인하여 종속변수에 대한 영향 정도를 확인하였다(지성호, 강영순, 2014). 구조방정식 모델을 통한 연구모형 분석의 결과는 [그림 4-1]과 같다.



[그림 4-1] 연구모형 검증 결과

*p<.05, **p<.01, ***p<.001

- a) 모든 수치는 표준화 계수 수치임
- b) 유의미한 경로는 실선으로 표시함

구조방정식 각 경로를 분석한 결과는 다음 [표 4-10]과 같다.

[표 4-10] 구조모형 분석 결과

가설	가설경로	Estimate		S.E.	C.R.	결과
		<i>B</i>	β			
H1	사용용이성 → 지각된 유용성	.795	.650	.076	10.481***	채택
H2	사용용이성 → 사용의도	.381	.180	.103	3.697***	채택
H3	사용용이성 → AI 리터러시	.242	.224	.070	3.484***	채택
H4	지각된 유용성 → AI 리터러시	.481	.543	.062	7.780***	채택
H5	지각된 유용성 → 사용의도	.759	.437	.100	7.609***	채택
H6	AI 리터러시 → 사용의도	.669	.340	.100	6.712***	채택

*p<.05, **p<.01, ***p<.001

구조모형 분석을 통한 가설 검증 결과를 살펴보면 다음과 같다.

첫째, ‘사용용이성’과 ‘지각된 유용성’ 간의 경로계수는 $\beta=.650(p<.001)$ 로 나타나 “사용용이성은 지각된 유용성에 정(+)의 영향을 미칠 것이다”라는 가설1(H1)은 채택되었다. AI 기술의 사용이 편리해짐에 따라 사용자의 업무 능력을 향상하고 유용성에 긍정적인 영향을 미친다.

둘째, ‘사용용이성’과 ‘사용의도’ 간의 경로계수는 $\beta=.180(p<.001)$ 로 나타나 “사용용이성은 사용의도에 정(+)의 영향을 미칠 것이다”라는 가설2(H2)는 채택되었다. 사용용이성은 사용의도에 영향을 미치는 것으로 나타났다. 이러한 결과는 AI 기술에 대한 사용용이성과 사용의도가 서로 독립적이지 않는 요소임을 시사한다. 즉, 사용자들이 AI 기술 사용이 편리함을 느낀다면, 이러한 편리함이 새로운 기술의 사용 의향이 있는 것임을 알 수 있다. 이진(2024)의 연구에서는 본 연구의 결과와는 반대로 사용용이성은 사용의도에 영향을 미치지 않는 것으로 확인되었다.

셋째, ‘사용용이성’과 ‘AI 리터러시’ 간의 경로계수는 $\beta=.224$ ($p<.001$)로 나타나 “사용용이성은 AI 리터러시 수준에 정(+)의 영향을 미칠 것이다”라는 가설3(H3)은 정(+)의 영향을 미치는 것으로 채택되었다. AI 기술 사용이 편리해짐에 따라 AI 리터러시 즉, 경험하고 활용하는 수준은 높아질 수 있다.

넷째, ‘지각된 유용성’과 ‘AI 리터러시’ 간의 경로계수는 $\beta=.543$ ($p<.001$)로 나타나 “지각된 유용성은 AI 리터러시에 정(+)의 영향을 미칠 것이다”라는 가설4(H4)는 채택되었다. AI 기술의 사용으로 인해 사용자의 업무 효율이 높아지고 미래의 삶에 영향을 미칠 수 있음에 따라 AI 리터러시 수준도 향상될 것으로 예상된다.

다섯째, ‘지각된 유용성’과 ‘사용의도’ 간의 경로계수는 $\beta=.437$ ($p<.001$)로 나타나 “지각된 유용성은 사용의도에 정(+)의 영향을 미칠 것이다”라는 가설5(H5)는 채택되었다. AI 기술의 유용하고 효과적임을 알게 됨에 따라 사용자의 사용의도는 높아질 것이다.

여섯째, ‘AI 리터러시’와 ‘사용의도’ 간의 경로계수는 $\beta=.340$ ($p<.001$)로 나타나 “AI 리터러시는 사용의도에 정(+)의 영향을 미칠 것이다”라는 가설6(H6)은 채택되었다. AI 기술에 대한 이해도가 높고 활용 역량이 많을수록 해당 기술을 더 많이 사용하고자 하는 의지가 높아지는 것으로 분석되었다.

3) 매개효과 검증

다음으로, 지각된 유용성과 사용의도 간 그리고 사용용이성과 사용의도 간 AI 리터러시가 매개작용하여 간접효과가 있는지 가설을 검증하였다. 매개변수 경로 간접효과 분석 결과는 아래의 [표 4-11]과 같다. 매개효과 유의성을 확인하기 위해 부트스트래핑(Bias Corrected Bootstrapping) 검정법을 실시하였으며, Estimate, 표준오차(S.E.), 부트스트랩(Bootstrap) 95% 신뢰구간 값을 분석하였다.

[표 4-11] 경로 간접효과

가설	변수	Estimate	S.E.	95% CI		결과
				Lower	Upper	
H7	지각된 유용성 → AI 리터러시 → 사용의도	.185	.034***	.118	.254	채택
H8	사용용이성 → AI 리터러시 → 사용의도	.481	.044***	.400	.570	채택

*p<.05, **p<.01, ***p<.001

지각된 유용성과 사용의도 간에 AI 리터러시의 매개효과는 95% 신뢰구간에서 .118~.254의 상한값과 하한값을 나타내고 있으며 0을 포함하지 않는 것으로 나타났다. 즉 AI 리터러시의 매개효과 유의성을 검증한 결과 간접효과가 있는 것으로 확인되었다($p<.001$). 사용용이성과 사용의도간 AI 리터러시의 매개효과는 95% 신뢰구간에서 .400~.570의 상한값과 하한값을 나타내고 있으며 0을 포함하지 않는 것으로 나타났다. 즉 AI 리터러시의 매개효과 유의성을 검증한 결과 간접효과가 있는 것으로 검증되었다($p<.001$).

지각된 유용성 및 사용용이성과 사용의도 간에 AI 리터러시의 매개효과가 검증됨에 따라 각 변수의 직접, 간접 및 총효과를 검증하였다[표 4-12].

[표 4-12] 표준화 직접효과, 간접효과 및 총효과

가설	경로	직접효과	간접효과	총효과	간접신뢰구간
H7	지각된 유용성 → AI 리터러시 → 사용의도	.437	.185	.622	.118~.254***
	지각된 유용성 → AI 리터러시	.543	.000	.543	.000~.000
	AI 리터러시 → 사용의도	.340	.000	.340	.000~.000
H8	사용용이성 → AI 리터러시 → 사용의도	.180	.481	.660	.400~.570***
	사용용이성 → AI 리터러시	.224	.353	.577	.000~.000
	AI 리터러시 → 사용의도	.340	.000	.340	.000~.000

*p<.05, **p<.01, ***p<.001

지각된 유용성과 사용의도 간 AI 리터러시가 매개작용을 하여 간접효과가 있으며, 각 변수 간 직접효과 .437, 간접효과 .185로 총효과는 .622로 나타났다. 사용용이성과 사용의도 간 AI 리터러시가 매개작용을 하여 간접효과가 있으며, 각 변수 간 직접효과 .180, 간접효과 .481로 총효과는 .660으로 나타났다. 즉, 지각된 유용성 및 사용용이성과 사용의도 간에 AI 리터러시가 각각 부분 매개 작용을 하는 것으로 확인되었다. 또한 AI 리터러시의 간접효과는 지각된 유용성 대비 사용용이성이 사용의도 간 매개 작용 시 상대적으로 효과가 큰 것으로 나타났다.

연구가설에 관한 연구 결과를 종합적으로 정리하면 다음 [표 4-13]과 같다.

[표 4-13] 연구가설 검정 결과 요약

가설번호	가설 내용	채택여부
기술수용모델(TAM) 검증		
H1	AI 기술의 사용용이성은 지각된 유용성에 정(+)의 영향을 미칠 것이다.	채택
H2	AI 기술의 사용용이성은 기술 사용자의 사용의도에 정(+)의 영향을 미칠 것이다.	채택
H3	AI 기술의 사용용이성은 AI 리터러시 수준에 정(+)의 영향을 미칠 것이다.	채택
H4	AI 기술의 지각된 유용성은 AI 리터러시 수준에 정(+)의 영향을 미칠 것이다.	채택
H5	AI 기술의 지각된 유용성은 AI 기술 사용의도에 정(+)의 영향을 미칠 것이다.	채택
H6	AI 리터러시 수준은 AI 기술 사용의도에 정(+)의 영향을 미칠 것이다.	채택
매개효과 검증		
H7	AI 리터러시 수준은 AI 기술의 지각된 유용성과 사용의도 간에 매개 작용을 할 것이다.	채택
H8	AI 리터러시 수준은 AI 기술의 사용용이성과 사용의도 간에 매개 작용을 할 것이다.	채택

제 5 장 결론

제 1 절 연구결과

본 연구는 새로운 기술과 서비스를 받아들이는 이용자들의 기술 수용 과정 및 사용의도를 파악하기 위해 기술수용모델을 기반으로 설명하였다. 또한 AI 기술에 대한 이해도, 활용 능력 및 비판 역량으로 대표되는 AI 리터러시 수준과 사용의도 간의 관계에 관해 설명하고 있다. 구체적으로 AI에 대한 기술수용모델의 변수로 지각된 유용성과 사용용이성을 설정하였고 AI 리터러시라는 요인을 매개변수로 설정하여 AI를 지속적으로 사용하려는 의도에 영향을 미칠 것인지 구조방정식 모델을 통해 검증하였다.

이 연구는 ChatGPT, 바드 및 빙과 같은 AI 기술이 일반 대중에게 널리 알려지고 사용되기 시작한 2023년에 사회적가치연구원에서 만 20세 이상 성인 1,000명을 대상으로 실시한 ‘2023 한국인이 바라본 사회문제’의 설문 조사 결과를 비교 및 참조하였으며, AI 사용률이 증가한 2024년 12월 온라인 설문조사를 통해 312명을 표본으로 수집된 설문 조사 결과를 기반으로 분석하였다. AI 기술에 대한 조사 데이터를 활용하여 실증분석을 수행하였다. 주요 연구 결과는 다음과 같다.

기술수용모델의 주요 변수 중 사용용이성은 지각된 유용성에 정(+)-적으로 영향을 미치고, 사용의도에도 유의미한 영향을 미치는 것으로 나타났다. AI 리터러시에는 가설과 같이 정(+)-적인 영향을 미치는 것으로 나타났다. 지각된 유용성은 AI 리터러시 수준과 사용의도에 정(+)-적인 영향을 미치는 것으로 나타났다. AI 리터러시 수준도 사용의도에 정(+)-적인 영향을 미치는 것으로 나타났다. 그리고 사용용이성과 지각된 유용성이 사용의도에 영향을 미치는 데 있어 AI 리터러시 수준이 매개 역할을 통한 간접효과를 일으키는지 확인 결과 지각된 유용성과 사용용이성이 유의미한 영향을 미치는 것으로 확인되었으며 사용용이성이 지각된 유용성 대비 총효과가 큰 것으로 확인되었다. 새로운 과학기술인 AI 기술을 도입하고

실행할 때, AI 리터러시 수준이 사용의도에 큰 영향을 미친다는 것을 확인하였다. 또한 새로운 기술에 대한 지각된 유용성이 향상되면 AI 기술에 대한 AI 리터러시 수준이 증가하고, 사용의도의 긍정적인 측면이 강화된다는 것을 실증적으로 검증하였다.

제 2 절 시사점

본 연구는 AI 기술의 사용의도에 영향을 미치는 요인 중 그 중요성이 대두되고 있는 AI 리터러시를 기존의 이론적 모델의 적용 및 확장을 시도하고 실증적으로 검증하였다는 점에서 의의가 있으며, 사용자가 AI 기술을 수용하는 데 AI 리터러시 수준이 매개효과가 있는 것을 밝혔다는 점에서 의의가 있다. 또한 향후 급속한 AI 기술의 발전에 따라 예상되는 AI 기술 사용자의 양극화 이슈 및 AI 리터러시 격차 해소를 위한 정책 수립 시 본 연구의 실증 결과를 기초 자료로 제공할 수 있다는 점에서 학문적 의의가 있다.

또한 본 연구의 분석 결과를 바탕으로 다음과 같은 정책적 시사점을 도출할 수 있다.

첫째, 사회적가치연구원의 설문조사에서 AI 기술 이용 경험과 이해도 즉 AI 리터러시 수준을 질문하고 사회 계층별로 분석한 결과, ‘수도권에 거주하고 40대 이상으로 고학력인 남성일수록 AI 활용 경험이 많고 유용하다’고 분석되었고 앞으로도 활용할 가능성도 높게 나타났다. 반면, 활용 경험이 적고 불편함을 느낀 사람들은 앞으로도 이용할 가능성이 크게 떨어지는 것으로 나타났다. 앞으로 AI 기술 발전이 지속됨에 따라 사회 계층별 새로운 형태의 AI 리터러시 격차는 지속적으로 커질 것으로 전망되며, 이를 적극적이고 선제적으로 해소하기 위한 고민이 필요하다.

선행연구에서 알려진 것처럼 디지털 리터러시 개념은 디지털 정보에 대한 이해도 증대와 디지털기기의 활용 역량에 초점이 맞추어져 있었다. 반면, AI는 방대하고 복잡한 데이터로부터 정보를 추출하고 이들 간의 상

호관계에서 의미를 찾아내어 추론하므로, 인간의 추론을 보완할 수 있는 기능을 할 수 있다(Androutsopoulou, Karacapilidis, Loukis, & Charalabidis, 2019). 즉, 정보를 이해하고 활용하는 데 있어 인간이 결정하는 영역과 중첩될 수 있다. 이것은 AI 리터러시 수준이 낮은 이용자라도 정보를 쉽게 이용하게 하는 측면이 있지만, 오히려 인간의 판단을 저해하여 AI에 의존하게 하는 요인이 될 수 있다(Barth & Arnold, 1999). 따라서 단순한 정보기기 활용 역량과 관련된 리터러시가 아닌 보다 포괄적 측면에서 비판적 리터러시 역량의 필요성이 요구된다(Eshet-Alkalai & Chajut, 2010). AI 리터러시 격차 해소를 위해서는 다양성, 형평성과 포용성(DE&I: Diversity, Equity, & Inclusion) 개념이 필수적으로 반영되어야 한다.

AI 기술의 발전과 적용에 있어서 다양성, 형평성과 포용성을 강화하기 위해서는 AI의 설계, 개발, 배포 과정에서 다양한 인종, 성별, 문화적 배경 등을 고려한 이해관계자들의 참여와 협력을 촉진하고, 소수집단의 목소리를 반영하는 것이 효과적이다(Fosch-Villaronga & Poulsen, 2022; 데이터동향, 2022). 현재 일부 기업들은 다양성을 증진하기 위해 AI 개발팀 내 다양한 배경을 가진 인재를 채용하고 있으며, 이를 통해 DE&I를 강화하는 노력을 기울이고 있다(데이터동향, 2022).

둘째, 현재 일상생활에서 남녀노소 불편함 없이 사용 중인 휴대전화처럼 AI 기술 사용을 확대하기 위해 AI의 사용용이성과 지각된 유용성에 관한 긍정적 인식 확산이 필요하다. 특히, 전체 국민 3명 중 2명은 고용 및 노동 불안정 문제를 AI 기술로 인해 가장 악화할 것으로 예상한다(사회적가치연구원, 2023). 부정적 인식의 극복을 위해 사회안전망의 합의와 구축, 노동자의 재교육 시스템 등 국가, 기업, 국민 간의 합의를 통한 적극적인 대응책이 요구된다.

현재의 AI 기술은 긍정적 및 부정적인 효과가 동시에 나타나고 있고, 이에 따라 AI의 추론 능력이 사람과 크게 차이를 나타내지 않는 분야 위주로 효율성과 생산성을 향상하는 방향으로 진행되고 있다(윤상오, 이은미, 성욱준, 2018). 챗봇 서비스와 같은 공공 민원 서비스를 도입하고 있는 기관에서는 AI 기술의 효능에 대해 긍정적인 효과를 인지하고는 있지만, 이

를 이용하는 개별 시민의 인식과는 격차가 확인되고 있다(김민진, 2021). 또한, AI의 오류로 인한 사고 및 부정확한 정보의 제공으로 인한 부정적 인식이 발생할 수 있으며 이에 관한 책임 소재의 제도적 보완이 요구되고 있다(김민진, 2021. 김성희, 2023).

이미 AI 기술이 우리 일상에 깊이 자리 잡고 있으며, 이를 사용하는 것이 불가피해지고 습관화되고 있다. 비록 생성형 AI의 인지적 개방성에 대해 위험을 인지하고 있지만, 서비스의 지속 사용 의도에는 영향을 미치지 않는다는 것을 확인하였다(장창기, 성옥준, 2022). 이는 사람들이 위험을 인지하고 있음에도 불구하고 지속적인 사용의도에는 영향을 미치지 않기 때문에, 시간이 지남에 따라 이에 따른 문제가 발생할 가능성이 존재한다는 것이다. 따라서 AI에 대한 기술적인 문제뿐만 아니라 법적, 윤리적 측면에서도 논의와 규제가 필요함을 시사한다.

셋째, 개인정보 유출 등의 문제로 인해 AI 기술에 대한 신뢰도가 떨어지는 경우가 있으며, 이로 인해 AI 기술의 활용도가 저하될 수 있다. 따라서 개인 정보 보호를 위한 기술적 대책 마련과 함께, AI의 오류를 최소화하고 신뢰도를 높이는 방안을 모색해야 한다. 개인 정보 유출에 관한 우려는 새로운 기술의 사용에 대한 불확실성을 확대한다. 개인 정보 유출에 관한 걱정과 정보 침해 경험은 AI와 같은 새로운 기술이 도입된 서비스 이용에 부정적인 영향을 미친다(김창일, 허덕원, 이혜민, 성옥준, 2019). AI 기술이 산업 전 분야에 도입되면서 방대한 자료에 접근할 수 있음에 따라 AI의 편견, 오류 및 사고 등이 발생하면 심각한 문제로 진행할 가능성이 크다(윤상오, 이은미, 성옥준, 2018). 따라서 개인, 기업, 정부 등 정보 관리의 주체들 모두가 AI의 오류를 줄이기 위한 대응 방안을 마련해야 한다. AI 기술에 대한 이해와 활용 능력을 높이기 위한 교육 및 기술 역량 강화 프로그램 개발과 지원이 필요하다(세계일보, 2023다). AI 알고리즘의 투명성을 높여 사용자들이 AI의 작동 방식을 이해할 수 있도록 하고 AI의 학습 데이터에 대한 검증과 평가를 강화하여 AI 기술의 성능과 정확도를 높여야 한다. 또한, AI 기술의 부작용을 예방하고 대처할 수 있는 법적 제도와 규제를 마련하여 사용자들의 권리와 안전을 보장할 필요가 있다.

제 3 절 한계 및 향후 연구

본 연구는 다음과 같은 한계가 존재한다. 첫째, 본 연구의 설문조사 데이터는 한국 내 사용자에게 한정되어 있으므로 다른 사회적 문화적 경제적 배경을 가진 지역의 사용자들에게는 직접적으로 적용될 수 없다는 한계가 있다. 둘째, 설문조사는 특정 시점에 이루어진 것이므로, 시간이 지남에 따라 변할 수 있는 기술수용성 소비자 인식과 행동 패턴을 완전히 반영하지 못할 수 있다. 특히, ChatGPT와 같은 생성형 AI는 단시간 내 급격한 기술 진보 과정을 나타내는바, 본 연구의 횡단면 연구로서의 한계가 있다. 셋째, 본 연구의 설문 문항과 데이터를 온라인 설문조사를 통해 제한된 계층 인원들을 표본으로 획득하였다. 이에 따라 AI 기술을 수용하는 사용자의 인식 수준을 측정할 수 있는 다양한 지표 중 저자가 채택한 지각된 유용성, 사용용이성, AI 리터러시 및 사용의도의 변수들이 전체 사용자의 인식을 대표하지 못할 수도 있다는 한계가 존재한다. 또한 실증분석을 실시할 때, 설문조사 데이터가 관측 변수와 잠재 변수를 모두 반영하지 못하였으며, 최근 주목받고 있는 AI 기술의 외생변수들인 법적, 윤리적 사회문제에 대한 변수를 충분히 고려하지 못하였다.

이러한 한계에도 불구하고 본 연구에서는 AI 기술 발전이 일반대중 및 산업계로 빠르게 확산하고 있는 도입기 시점에 민간과 공공기관에서 과제 연구 및 정책 실행 시 참고할 수 있는 기초 자료를 제공했다는 점에서 학문적 시사점을 찾을 수 있다. 실증분석 결과를 토대로 앞으로는 더 많은 대상 집단과 다양한 요인들을 포함한 연구를 진행함으로써 보다 정확하고 신뢰도가 높은 결과를 도출할 필요가 있다. 이러한 노력을 통해 AI 기술의 발전과 대중화에 이바지할 수 있을 것이다. 특히 AI 리터러시의 경우 학력에 따라 격차가 크게 발생하는 것으로 나타나므로, 고학력자가 대부분의 인적 구성원인 연구 기관에서 AI 기술을 사용, 개발을 위해 활용할 때는 AI에 대한 지각된 유용성, 사용용이성, AI 리터러시를 고려해야 한다. 이를 통해 보다 효과적인 AI 기술 적용 방안을 모색할 수 있으며, AI 기술의 대중화와 발전에도 이바지할 수 있을 것이다. 또한, 성별이나 연령

대 등 다른 요인들도 함께 고려하여 종합적인 분석을 실시하면 더욱 정확한 결과를 얻을 수 있을 것이다. 앞으로는 이와 같은 다양한 요인들을 고려한 연구가 지속적으로 이루어져야 하며, 이를 통해 AI 기술의 발전과 대중화에 이바지할 수 있도록 해야 한다. 향후 AI 기술 관련 연구 방향으로 AI 기술에 대한 법적·윤리적 사회문제를 고려한 AI 리터러시 개념 정의, 역량 강화 및 활용 방안에 대한 논의를 발전시키는 것이 필요하다.

참 고 문 헌

1. 국내문헌

- 고영화, 임춘성. (2021). 인공지능 기술수용과 윤리성 인식이 이용의도에 미치는 영향. 『디지털융복합연구』, 19(3), 217-225.
- 고윤정. (2022). AI 리터러시 향상이 자기효능감에 미치는 영향. 『한국지식정보기술학회 논문지』, 17(5), 1017-1028.
- 곽준도, 전종우. (2023). 종합적 사고경향, 지각된 유용성, 인공지능 윤리적 위험 인식이 인공지능 기술의 태도와 사용 의도에 미치는 영향. 『문화와 융합』, 45(5), 521-536.
- 권성호, 김성미. (2011). 소셜 미디어 시대의 디지털 리터러시 재개념화. 『미디어와 교육』, 1(1), 65-82.
- 권순동. (2015). SEM에서 위계모형을 이용한 다중공선성 문제 극복방안 연구: 소셜커머스의 재구매의도 영향요인을 중심으로. 『한국정보기술응용학회』, 22(2), 149-169.
- 김계수. (2010). 『AMOS 18.0 구조방정식모형 분석』. 서울: 한나래아카데미.
- 김민진. (2021). 공공안전부문 인공지능(AI) 도입 현황 및 시사점. 『한국통신정책연구원: AI Trend Watch』, 2021(17), 1-10
- 김성주. (2021). 인공지능 리터러시 향상을 위한 앱 개발 초등 교육 프로그램. 서울교육대학교 교육전문대학원 석사학위논문, 1-66
- 김성희. (2023). ChatGPT에 대한 의인화 지각이 인지된 혜택과 인지된 위험을 매개로 지속 사용 의도에 미치는 영향: AI 리터러시의 조절된 매개효과. 동국대학교 대학원 석사학위논문.
- 김윤경. (2022). 인공지능 챗봇 서비스의 수용태도에 미치는 영향요인 분석: 서비스 가치 매개효과 중심으로. 『한국콘텐츠학회논문지』, 22(2), 255-269.

- 김주은. (2019). 인공지능이 인간사회에 미치는 영향에 대한 연구. 『국제문화기술진흥원』, 5(2), 255-269.
- 김진석. (2021). 인공지능 리터러시 기반 초·중등교육의 내용과 교수·학습 방안 탐구. 『한국초등교육』, 32(3), 19-35.
- 김창일, 허덕원, 이혜민, 성욱준. (2019). 침해 경험 및 정보보호 인식이 정보보호 행동에 미치는 영향에 대한 연구: 이중 프로세스 이론을 중심으로. 『정보화정책』, 26(2), 62-80.
- 김태령, 류미영, 한선관. (2020). 초중등 인공지능 교육을 위한 프레임워크 기초 연구. 『인공지능교육학회 인공지능연구 논문지』, 1(1), 31-42.
- 김태문, 한진수. (2009). 인터넷 여행상품의 고객구매의도에 관한 연구. 『관광연구』, 24(1), 185-204.
- 김태창, 변순용. (2021). AI 윤리교육의 필요성과 내용 구성에 관한 연구. 『인공지능인문학연구』, 8, 71-104.
- 김효정. (2018). 결혼이민자 여성소비자의 디지털정보 격차지수 결정요인: 연령별 차이 연구. 『대한가정학회』, 56(3), 217-232.
- 남일규, 김승권, 김준연, 김영종, 박희준. (2017). 기술수용모델을 기반으로 한 소프트웨어 품질격차 요인 연구. 『정보기술아키텍처연구』, 14(1), 45-54.
- 노경섭. (2019). 『제대로 알고 쓰는 논문 통계분석: SPSS & AMOS (개정증보판)』. 서울: 한빛아카데미.
- 데이터동향. (2022). 데이터 리터러시, AI, 스토리텔링, 형평성, 문화 및 관리. <https://www.tableau.com/ko-kr/reports/data-trends/> (최종접속일: 2024.01.28.).
- 머니투데이. (2023). 대화형 AI 사용해봤다 51.5%. 연령 직업이 가른 챗GPT시대. <https://news.mt.co.kr/mtview.php?no=2023091816513792609/> (최종접속일: 2024.01.28.).
- 문영임, 이성규, 김지혜. (2021). 장애인의 디지털정보화 활용 수준이 삶의 만족도에 미치는 영향: 사회적 지지의 조절효과 분석. 『정보화정책』, 28(4),

36-53.

박유진, 이상원. (2022). 고객 서비스 챗봇의 사회적 실재감이 서비스 품질과 사용자 인식에 미치는 영향. 『고객만족경영연구』, 24(1), 85-104.

박태웅. (2023). 『박태웅의 AI 강의』. 서울: 한빛비즈.

백민소, 신준섭, 신유선. (2021). 독거노인의 ICT 기반 돌봄 보조 기기 사용의향 및 필요 기능 인식에 대한 기술적 연구. 『한국노년학』, 41(1), 25-48.

4차산업혁명위원회. (2021). 4차위 대국민 조사결과, AI 시대 도래에 공감하고 AI 대중화 요구. <https://blog.naver.com/kr4thir/222491578814/> (최종접속일: 2024.01.28.).

사회적가치연구원. (2023). 『2023 한국인이 바라본 사회문제』. 서울: 사회적가치연구원(CSES).

삼일회계법인 PwC. (2024). 『PwC's 27th Annual Global CEO Survey, 끊임없는 혁신의 시대에서 성공하기』. 서울: pwc 삼일회계법인.

성동규. (2009). 중간광고에 대한 인지된 유용성 및 인지된 위험이 중간광고 허용 의사에 미치는 영향에 관한 연구. 『한국언론학보』, 53(6), 374-404.

성욱준, 황성수. (2017). 지능정보시대의 전망과 정책대응 방향 모색. 『정보화정책』, 24(2), 3-19.

송선영. (2021). 인공지능 윤리와 (AI) 디지털 리터러시 역량 함양 방안에 대한 토론편. 한국초등도덕교육학회 학술대회, 231-232.

송효진. (2014). 질적 정보격차와 인터넷 정보이용의 영향 요인 고찰: 이용자의 디지털 리터러시, 인식, 자기효능감을 중심으로. 『한국정책과학학회보』, 18(2), 85-116.

세계일보. (2023가). 편향된 '알고리즘'에 갇힌 세상... 'AI 문해력' 교육 급선무. <https://www.segye.com/view/20230907515674/> (최종접속일: 2024.01.28.).

세계일보. (2023나). AI와 공존하려면. <https://www.segye.com/newsView/20230>

908513363 / (최종접속일: 2024.01.28.).

세계일보. (2023다). 윤리적 사용 강조 이상으로 AI 기술 자체 이해 중요. <https://www.segye.com/view/20230907515668/>(최종접속일: 2024.01.28.).

신건권. (2013). 『석·박사학위 및 학술논문 작성 중심의 AMOS 20 통계분석 따라하기』. 서울: 청람.

신소영, 이승희. (2019). 디지털 리터러시 측정도구 개발 및 타당화 연구. 『학습자중심교과교육연구』, 19(7), 749-768.

심태용, 윤성준. (2020) 온라인 쇼핑몰 특성이 감성적 반응과 지각된 가치, 재이용의도에 미치는 영향: 확장된 기술수용모델(TAM2)을 중심으로. 『한국산학기술학회논문지』, 21(4), 374-383.

안승규, 고상훈, 최혁진, 황순현, 곽기영, 안현철. (2023). ChatGPT 사용자의 성격 특성이 만족 및 지속사용의도에 미치는 영향에 대한 분석. 『한국경영정보학회 학술대회논문집』, 2023(06), 628-629.

안정임, 서윤경. (2014). 디지털 미디어 리터러시 격차의 세부요인 분석 - 세대와 경제수준을 중심으로. 『디지털융복합연구』, 12(2), 69-78.

연합뉴스. (2021). 성희롱·혐오논란에 3주만에 멈춘 '이루다'... AI 윤리 숙제 남기다. <https://www.yna.co.kr/view/AKR20210111155153017> / (최종접속일: 2024.01.28.).

오설미, 최송식. (2021). 노인의 디지털 정보수준이 신기술 이용의사에 미치는 영향: 기술적 자기효능감과 이용성과의 다중매개효과를 중심으로. 『노인복지연구』, 76(4), 137-170.

오종철, 윤성준, 우원. (2010). 모바일 인터넷 서비스 이용의도에 관한 연구: 개정된 TRAM 모형을 중심으로. 『서비스경영학회지』, 11(5), 127-148.

우종필. (2012). 『구조방정식모델 개념과 이해』. 서울: 한나래출판사.

윤상오. (2018). 인공지능 기반 공공서비스의 주요 쟁점에 관한 연구. 『한국공공관리학보』, 32(2), 83-104.

- 윤상오, 이은미, 성욱준. (2018). 인공지능을 활용한 정책결정의 유형과 쟁점에 관한 시론. 『한국지역정보학회지』, 21(1), 31-59.
- 윤종현, 윤한성. (2024). 인공지능 리터러시 수준이 사용의도에 미치는 영향: 지각된 인식의 매개효과를 중심으로. 『혁신클러스터연구』, 14(2), 45-79.
- 윤태환, 왕새롬. (2023). Examining users' intention to use generative artificial intelligence for travel information search : An exploratory study of ChatGPT. 『관광연구논총』, 35(4), 197-221.
- 이강호. (2023). 대화형 AI 관련 인식도 조사 결과 및 정책적 시사점. <https://brunch.co.kr/@klee4u/28> / (최종접속일: 2024.01.28.).
- 이기호. (2017). 보건복지 부문 정보통신기술(ICT) 정책 추진 현황과 과제. 『보건복지포럼』, 250(-), 42-56.
- 이다점, 김성원, 이영준. (2021). 인공지능 리터러시 교육 연구 동향 분석. 『한국컴퓨터교육학회 학술발표대회논문집』, 25(2A), 25-27.
- 이데일리. (2024). 인간처럼 보고 듣고 말한다. 오픈 AI, GPT-4o 출시. <https://www.edaily.co.kr/News/Read?newsId=01325126638888920&mediaCodeNo=257&OutLnkChk=Y> / (최종접속일: 2024.11.13.).
- 이상은, 김예진, 구민영, 이수민. (2023). 1학기 vs. 2학기 생성형 AI 활용 학습경험 비교. 성균관대 교육개발센터. <https://chatgpt.skku.edu/chatgpt/notice.do?mode=view&articleNo=165899> / (최종접속일: 2024.11.10.).
- 이우진, 백혜진. (2022). 키워드 네트워크 분석을 통한 리터러시 교육 연구 동향. 『디지털융복합연구』, 20(5), 53-59.
- 이유미. (2022). 디지털 시대 새로운 패러다임과 리터러시 : 디지털 리터러시와 AI 리터러시를 중심으로. 『교양학연구』, -(20), 35-60.
- 이은창. (2023). 『인공지능과 사회문제 : 세 가지 차원의 위협』. 서울: 사회적가치연구원(CSES).
- 이지연, 이재익, 전지호. (2017). ICT를 활용한 독거노인 헬스케어서비스 실증 연구: 송파구 지역을 중심으로. 『디자인융복합학회』, 16(1), 257-280.

- 이철승, 백혜진. (2022). 디지털 리터러시 함양을 위한 교수학습 방법 연구. 『디지털융복합연구』, 20(5), 351-356.
- 이철승, 백혜진. (2023). 인공지능 챗봇 발전에 따른 AI 리터러시 필요성 연구. 『한국전자통신학회 논문지』, 18(3), 421-426.
- 이철현. (2020). AI 시대 역량 함양을 위한 실과 소프트웨어교육의 방향. 『實科教育研究』, 26(2), 41-64.
- 이한신, 김판수. (2019). 소비자의 기술수용과 저항이 인공지능(AI) 사용의도에 미치는 영향. 『經營學研究』, 48(5), 1195-1219.
- 임성진. (2012). 스마트폰 사용자요인이 업무성과에 미치는 영향에 관한 연구: 스마트폰 수용요인의 매개효과 검증을 중심으로. 한성대학교 지식서비스 & 컨설팅대학원 석사학위논문.
- 임성진, 한경석, 정미라. (2017). 공공기관 근무자의 스마트 모바일기기 사용과 업무성과의 관계에 관한 연구: TAM 모형을 활용한 업무성과의 관계 검증을 중심으로. 『한국디지털콘텐츠학회』, 18(7), 1465-1474.
- 장창기, 성옥준. (2022). 인공지능 기반 공공서비스 정책수용 의도에 관한 연구: 개인의 인식과 디지털 리터러시 수준이 미치는 영향을 중심으로. 『정보화정책』, 29(1), 60-83.
- 장한진, 노기영. (2017). 기술수용모델을 이용한 초기 이용자들의 가상현실기기 채택 행동 연구. 『디지털융복합연구』, 15(5), 353-361.
- 조민호. (2021). 인공지능의 역사, 분류 그리고 발전 방향에 관한 연구. 『한국전자통신학회 논문지』, 16(2), 307-312.
- 조진완, 이종호. (2014). 소비자의 기술준비도가 사물인터넷 사용의도에 미치는 영향에 관한 연구. 『한국경영교육학회 학술발표대회논문집』, 2014.12., 533-554.
- 주간동아. (2023). 김정한 대표 미국 주식 하나 고르라면 단연 엔비디아. <https://www.donga.com/news/Economy/article/all/20231209/122551336/1/> (최종접속일: 2024.11.10.).

- 주경희, 김동심, 김주현. (2018). 노년층의 정보격차에 대한 성별에 따른 차이분석과 예측변인 탐색. 『한국노인복지학회 학술대회』, 443-463.
- 중앙일보. (2023). 구글, 3만명 구조조정 추진설...커지는 AI發 해고 공포. <https://www.joongang.co.kr/article/25217626#home> / (최종접속일: 2024.01.28.).
- 지성호, 강영순. (2014). 사회과학분야의 구조방정식모형에서 매개효과 검정 방법에 대한 논의. 『한국자료분석학회』, 16(6), 3121-3131.
- 최재현. (2023). 『AI 리터러시』. 서울: 열린 인공지능.
- 파이낸셜뉴스. (2024). GPT 스토어 오픈...韓기업 입점 기대감속 종속우려 목소리도. <https://www.fnnews.com/news/202401111351273132> / (최종접속일: 2024.01.28.).
- 한국일보. (2024). 외국인은 못 읽던 '외계 한국어', 이제 안 통한다... 오픈AI, 추론하는 '괴물 AI' 공개. <https://www.hankookilbo.com/News/Read/A2024091313200001094> / (최종접속일: 2024.11.17.).
- 한정선, 오정숙. (2006). (21세기 지식 정보 역량 활성화를 위한) 디지털 리터러시의 조작적 정의 및 하위 영역 규명. 서울: 한국교육학술정보원.
- 황현정, 박지수, 김승완. (2023). 국내외 AI 리터러시의 연구 동향 분석: 체계적 문헌고찰 방법을 중심으로. 『교육문화연구』, 29(3), 135-164.
- 헤럴드경제. (2024). AI판 구글 앱스토어 되나. 오픈AI, 생태계 장악 우려. <http://biz.heraldcorp.com/view.php?ud=20240112000305> / (최종접속일: 2024.01.28.).
- 히든그레이스. (2023). 『한번에 통과하는 논문: AMOS 구조방정식 활용과 SPSS 고급분석』. 서울: 한빛아카데미.

2. 국외문헌

- Ajzen, I. (1991). The Theory of Planned Behavior. *Organizational Behavior and Human Decision Processes*, 50(2), 179–211.
- Ajzen, I. & Fishbein, M. (1975). Belief, attitude, intention and behavior: An introduction to theory and research. *Reading, MA: Addison-Wesley*, 578.
- Ajzen, I. & Fishbein, M. (1980). Theory of reasoned action as applied to moral behavior: A confirmatory analysis”. *Journal of Personality and Social Psychology*, 62(1), 98–109.
- Anderson, J. C. & Gerbing, D. W. (1988). Structural Equation Modeling in Practice: A review and recommended two-step approach. *Psychological Bulletin*, 103(3), 411–423.
- Androutsopoulou, A., Karacapilidis, N., Loukis, E., & Charalabidis, Y. (2019). Transforming the Communication between Citizens and Government through AI-guided Chatbots. *Government Information Quarterly*, 36, 358–367.
- Anthes, G. (2017). Artificial Intelligence Poised to Ride a New Wave, *Communications of the ACM*, 60(7), 19–21.
- Associated Press. (2023). In Hollywood writers’ battle against AI, humans win(for now). <https://apnews.com/article/hollywood-ai-strike-wga-artificial-intelligence-39ab72582c3a15f77510c9c30a45ffc8/> (Final access date: 2024.11.10.).
- Ayoub, K. & Payne, K. (2016). Strategy in the age of artificial intelligence. *Journal of strategic studies*, 39(5–6), 793–819.
- Bagozzi, R. P. & Yi, Y. (1988). On the evaluation of structural equation models. *Journal of the academy of marketing science*, 16, 74–94.

- Barth, T. J. & Arnold, E. (1999). Artificial intelligence and administrative discretion: Implications for public administration. *The American Review of Public Administration*, 29(4), 332–351.
- Bhattacharjee, A. (2001). Understanding Information Systems Continuance: An Expectation Confirmation Model. *MIS Quarterly*, 25(3), 351–370.
- Bloomberg. (2024). OpenAI Scale Ranks Progress Toward ‘Human-Level’ Problem Solving. <https://www.bloomberg.com/news/articles/2024-07-11/openai-sets-levels-to-track-progress-toward-superintelligent-ai/> (Final access date: 2024.11.16.).
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340.
- Davis, F. D., Bagozzi, R. P., & Warshaw, P. R. (1989). User acceptance of computer technology – A comparison of two theoretical models. *Management Science*, 35(8), 982–1003.
- DeLone, W. H. & McLean, E. R. (1992). Information Systems Success: The Quest for the Dependent Variable. *Information Systems Research*, 3(1), 60–95.
- de Sousa, W. G., de Melo, E. R. P., Bermejo, P. H. D. S., Farias, R. A. S., & Gomes, A. O. (2019). How and Where is Artificial Intelligence in the Public Sector Going? A Literature Review and Research Agenda. *Government Information Quarterly*, 36(4), 101392, 1–14.
- Elsholz, E., Chamberlain, J., & Kruschwitz, U. (2019). Exploring Language Style in Chatbots to increase Perceived Product Value and User Engagement.” In *Proc. 2019 Conf. Human Information Interaction and Retrieval* : 301–305.

- Emwanta, M. & Nwalo, K. I. N. (2013). Influence of computer literacy and subject background on use of electronic resources by undergraduate students in universities in south-western Nigeria. *International Journal of Library and Information Science*, 5(2), 29–42.
- Eshet-Alkalai, Y. & Chajut, E. (2010). You Can Teach Old Dogs New Tricks: The Factors That Affect Changes over Time in Digital Literacy. *Journal of Information Technology Education: Research*, 9(1), 173–181.
- Euractiv. (2023). OECD updates definition of Artificial Intelligence ‘to inform EU’s AI act’. <https://www.euractiv.com/section/artificial-intelligence/news/oecd-updates-definition-of-artificial-intelligence-to-inform-eus-ai-act/> (Final access date: 2024.01.28.).
- Fornell, C. & Larcker, D. F. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of marketing research*, 18(1), 39–50.
- Fosch-Villaronga, E. & Poulsen, A. (2022). Diversity and inclusion in artificial intelligence. *Law and Artificial Intelligence: Regulating AI and Applying AI in Legal Practice*, 109–134.
- Haenlein, M. & Kaplan, A. (2019). A brief history of artificial intelligence: On the past, present, and future of artificial intelligence. *California management review*, 61(4), 5–14.
- Hargittai, E. & Hinnant, A. (2008). Digital inequality: Difference in young adults’ use of the internet.” *Communication Research*, 35(5), 602–621.
- IMF. (2024). Gen-AI: Artificial intelligence and the future of work. <https://www.imf.org/en/Publications/Staff-Discussion-Notes/Issues/2024/01/14/Gen-AI-Artificial-Intelligence-and-the-Future-of>

–Work–542379 (Final access date: 2024.01.28.).

- Jenkins, H. (2009). *Confronting the challenges of participatory culture: Media education for the 21st century*. London, England: The MIT Press.
- Jackson, C. M., Chow, S., & Leitch, R. A. (1997). Toward an Understanding of the Behavioral Intention to Use an Information System. *Decision Sciences*, 28(2), 357–389.
- Kim, J. H. & Nam, C. G. (2019). *Analyzing continuance intention of recommendation algorithms. 30th European Conference of the International Telecommunications Society (ITS : International Telecommunications Society): Towards a Connected and Automated Society*, Helsinki, Finland, 16th–19th June, 2019.
- Kong, S. C. (2008). A curriculum framework for implementing information technology in school education for fostering information literacy. *Computers & Education*, 51(1), 129–141.
- Long, D. & Magerko, B. (2020). *What is AI literacy? competencies and design considerations*. Paper presented at the Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems, 1–16.
- Lutz, C. (2019). Digital inequalities in the age of artificial intelligence and big data. *Human Behavior and Emerging Technologies*, 1(2), 141–148.
- Molnar, S. (2002). Explanation frame of the digital divide Issue. *Information Society*, 4, 102–118.
- Morris, M. R., Sohl-dickstein, J., Fiedel N., Warkentin, T., Dafoe, A., Faust, A., Farabet, C., & Legg, S. (2023). Levels of AGI: Operationalizing progress on the path to AGI. *arXiv: 2311.02462*, 1–19.

- Mueller, R. O. & Hancock, G. R. (2010). Structural Equation Modeling. In G. R. Hancock & R. O. Mueller (Eds.), *The Reviewer's Guide to Quantitative Methods in the Social Sciences* : 371–383, New York: Routledge.
- Netemeyer, R. G., Boles, J. S., McKee, D. O., & McMurrian, R. (1997). An Investigation into the Antecedents of Organizational Citizenship Behaviors in a Personal Selling Context, *Journal of Marketing*, 61(3), 85–98.
- Pavlou, P. A. (2003). Consumer Acceptance of Electronic Commerce: Integrating Trust and Risk with the Technology Acceptance Model. *International Journal of Electronic Commerce*, 7(3), 101–134.
- Rai, A., Lang, S. S., & Welker, R. B. (2002). Accessing the Validity of IS Success Models: An Empirical Test and Theoretical Analysis. *Information systems research*, 13(1), 50–69.
- Rogers, E. (2003). *Diffusion of Innovations 5th ed.*. New York, NY: Free Press, 196.
- Russell, S. J. & Norvig, P. (2010). *Artificial intelligence a modern approach*. London.
- Ruy, H. I. & Jo J. W. (2021). Development of an educational system for artificial intelligence education for K–12 targets based on 4 P. *Digital Convergence Research*, 19(1), 141–149.
- Schank, R. C. (1991). Where's the AI?. *AI magazine*, 12(4) : 38–49.
- Segars, A. H. (1997). Assessing the Unidimensionality of Measurement: A Paradigm and Illustration within the Context of Information Systems Research. *Omega*, 25(1), 107–121.
- Sharma, G. D., Yadav, A., & Chopra, R. (2020). Artificial Intelligence and Effective Governance: A Review, Critique and Research Agen

- da. *Sustainable Futures*, 2, 100004.
- Similarweb. (2024). OpenAI. <https://www.similarweb.com/website/openai.com/#overview> (Final access date: 2024.01.28.).
- Shin, D. H. (2009). An empirical investigation of a modified technology acceptance model of IPTV. *Behaviour & Information Technology*, 28(4), 361–372
- Venkatesh, V. (2000). Determinants of Perceived ease of Use: Integrating control, intrinsic motivation, and emotion into the technology acceptance model. *Information Systems Research*, 11(4), 342–365.
- Venkatesh, V. & Davis, F. D. (1996). A model of the antecedents of perceived ease of use: Development and test. *Decision sciences*, 27(3), 451–481.
- Venkatesh, V. & Morris. M. G. (2000). Why don't men ever stop to ask for directions? Gender, social influence, and their role in technology acceptance and usage behavior. *MIS Quarterly*, 24(2), 115–139.
- Venkatesh, V. & Morris. M. G., & Davis, F. D. (2003). User Acceptance of Information Technology: Toward a Unified View. *MIS Quarterly*, 27(3), 425–478.
- Wang, Y. & Hou, J. (2019). A Study on the Choice Preference of Man-made/Robot Service Mode. *Shanghai Management Science*, 41(6), 37–42.

[부 록]

기술수용모델을 활용한 AI 기술이 사용의도에 미치는 영향에 관한 연구 설문지

안녕하십니까.

바쁘신 와중에 본 설문을 위해 시간을 내어 주셔서 감사드립니다.

본 설문조사는 AI 기술에 대한 사용자들의 인식을 조사하여 AI 기술이 사용자들의 기술 수용에 미치는 요인에 관한 연구를 위해 기초자료를 수집하기 위한 것입니다. 본 설문조사는 대략 10분 내외의 시간이 소요될 것으로 예상됩니다.

본 설문에는 맞거나 틀리는 내용이 없으며, 귀하의 의견은 통계법 33, 34조에 의거 통계적 기초 자료로서 오직 연구 목적으로만 사용되며 철저한 보안이 유지됨을 강조 드립니다. 본 연구에 대해서 의문이 있으시거나, 연구결과에 관심이 있으신 분은 아래의 연락처로 문의하여 주십시오. 바쁘시더라도 평소의 생각이나 의견을 솔직하고 성실하게 답변해주실 것을 부탁드립니다.

설문조사에 응해 주셔서 감사합니다.

소속학교 및 학과 : 한성대학교 지식서비스&컨설팅 대학원,

지식서비스&컨설팅 학과, ESG경영컨설팅 전공

연구자 : 윤종현 (arnoldyoon71@naver.com, 010-8293-0363)

지도교수 : 정진택

설문일시 : 2024. 11. 28.

AI 기술 (생성형 AI : Generative AI) 이란?

사용자의 특정한 요구사항에 따라 결과물을 능동적으로 생성하는 인공지능(AI) 기술을 의미하며, 입력된 데이터를 학습하고, 학습된 데이터를 기반으로 새로운 데이터 결과를 만들어내는 기술을 의미합니다.

AI 기술(생성형 AI) 종류는?

ChatGPT(OpenAI), 바드(Bard, 구글), Copilot(마이크로소프트), Claude(Anthropic), Cue(네이버), 뽀빠, 미드저니(Midjourney), DALL-E, Perplexity 등

AI 리터러시란?

AI 기술과 각종 응용 프로그램을 이해, 활용하는 데 필요한 지식, 기술 및 윤리적 인식을 포함하는 포괄적 개념을 의미하며, 기술적 이해를 넘어 비판적 사고, 윤리적 고려 및 AI의 광범위한 사회적 영향에 대한 인식을 포함하는 다차원적 개념을 의미합니다.

Q1. 귀하는 ChatGPT(OpenAI), 바드(Bard, 구글), Copilot(마이크로소프트), Claude (Anthropic), Cue(네이버) 등과 같은 AI 기술 (생성형 AI)를 사용해 본 경험이 있으십니까?

① 있다 ② 없다

Q2. (지각된 사용용이성) 다음은 AI 기술(생성형 AI)의 이용 경험에 대한 질문입니다. 각 문항에 대하여 해당되는 곳에 표시해 주시기 바랍니다.

구분	전혀 그렇지 않다	그렇지 않다	보통 이다	그렇다	매우 그렇다
1. AI의 이용 방법은 쉽다고 생각한다					
2. AI 이용에 금방 익숙해질 수 있다고 생각한다					
3. AI의 방법은 명확하고 간단하다고 생각한다					
4. AI는 내가 원하는 것을 쉽게 할 수 있게끔 해준다					
5. AI를 사용하면 문제 해결을 쉽게 할 수 있다					

Q3. (지각된 유용성) 다음은 AI 기술(생성형 AI)의 이용 경험에 대한 질문입니다. 각 문항에 대하여 해당되는 곳에 표시해 주시기 바랍니다.

구분	전혀 그렇지 않다	그렇지 않다	보통 이다	그렇다	매우 그렇다
1. AI를 이용하면 내 목적을 더 빠르게 달성할 수 있다고 생각한다					
2. AI를 이용하면 내 관심사나 업무를 효율적으로 처리하는데 도움이 된다					
3. AI를 사용하면 유용한 정보를 얻을 수 있다고 생각한다					
4. AI를 사용하면 생산성이 증가한다					
5. AI는 전반적으로 나에게 유용하다고 생각한다					

Q4. (AI 리터러시) 다음은 귀하의 정보 처리 소양에 대한 질문입니다. 각 문항에 대하여 해당되는 곳에 표시해 주시기 바랍니다.

구분	전혀 그렇지 않다	그렇지 않다	보통 이다	그렇다	매우 그렇다
1. 나는 AI가 응답한 정보를 이해할 수 있다					
2. 나는 AI가 무엇인지 설명할 수 있다					
3. 나는 어떻게 AI를 이용해 내게 도움이 되는 정보를 요구할 수 있는지 알고 있다					
4. 나는 AI가 응답한 정보가 믿을 만 한지 않은지 판단할 수 있다					
5. 나는 AI를 통해 얻은 정보가 나의 상황에 적용가능한지 선별할 수 있다					
6. 나는 AI를 통해 목적에 맞는 정보를 수집하고 판단할 수 있다					
7. 나는 AI를 통해 얻은 정보를 필요한 순간에 적절하게 활용할 수 있다					
8. 나는 AI를 통해 나에게 가장 도움이 되는 정보가 무엇인지 구분할 수 있다					
9. 나는 AI를 통해 빠른 시간 내에 필요한 정보를 찾을 수 있다					
10. 나는 AI를 통해 다양한 생각과 의견을 검색할 수 있다					
11. 나는 AI 결과물에 대해 그대로 받아들이는 것이 아니라 해석이 필요하다는 것을 이해하고 있다					
12. 나는 AI가 응답한 정보가 거짓이거나 오류일 수 있다는 것을 이해하고 있다					

Q5. (일자리 영향) 다음은 AI가 미치게 될 일자리 영향에 대한 질문입니다. 각 문항에 대하여 해당되는 곳에 표시해 주시기 바랍니다.

구분	전혀 그렇지 않다	그렇지 않다	보통 이다	그렇다	매우 그렇다
1. AI는 나의 일자리(직업)에 위협 (위험)이다					
2. AI로 인해 나의 수입이 감소할 것이다					
3. AI로 인해 나의 경력(Career) 발전에 부정적인 영향을 미칠 것이다					
4. AI로 인해 나의 사회적 지위나 인식에 부정적 영향을 미칠 것 이다					
5. AI로 인해 나는 재교육 및 훈련 이 필요할 것이다					
6. AI로 인해 나는 직업전환이나 업 무전환이 필요할 것이다					

Q6. (긍정 영향) 다음은 AI가 미치게 될 영향에 대한 귀하의 인식에 대한 질문입니다. 각 문항에 대하여 해당되는 곳에 표시해 주시기 바랍니다.

구분	전혀 그렇지 않다	그렇지 않다	보통 이다	그렇다	매우 그렇다
1. AI가 내 삶의 질을 향상시켜줄 것이다					
2. AI가 국가, 기업, 개인의 경제적 성장과 발전에 기여할 것이다(비 용절감, 새로운소득 등)					
3. AI가 교육과 학습을 개선할 것 이다					
4. AI가 건강과 의료 분야를 개선 할 것이다					
5. AI가 환경 보호와 지속가능한 발전에 도움을 줄 것이다					
6. AI가 우리 사회의 다양성, 형평성 과 포용성을 개선시킬 것이					

Q7. (부정 영향) 다음은 AI가 미치게 될 영향에 대한 귀하의 인식에 대한 질문입니다. 각 문항에 대하여 해당되는 곳에 표시해 주시기 바랍니다.

구분	전혀 그렇지 않다	그렇지 않다	보통 이다	그렇다	매우 그렇다
1. AI가 개인정보 침해나 프라이버시 침해 우려를 키우는 데 영향을 미칠 것이다					
2. AI가 허위조작정보(가짜뉴스)나 딥페이크 생성에 영향을 미칠 것이다					
3. AI가 윤리적 문제를 발생시킬 것이다					
4. AI가 인간의 자유와 권리를 침해할 것이다					
5. AI가 군사적 위협을 발생시킬 것이다					
6. AI가 인간의 창의성과 독창성을 제한할 것이다					

Q8. (지속사용의도) 다음은 향후 AI 기술(생성형 AI)의 사용 의도에 대해 귀하의 생각을 묻는 질문입니다. 각 문항에 대하여 해당되는 곳에 표시해 주시기 바랍니다.

구분	전혀 그렇지 않다	그렇지 않다	보통 이다	그렇다	매우 그렇다
1. AI를 이용하면 내 목적을 더 빠르게 달성할 수 있다고 생각한다					
2. AI를 이용하면 내 관심사나 업무를 효율적으로 처리하는데 도움이 된다					
3. AI를 사용하면 유용한 정보를 얻을 수 있다고 생각한다					
4. AI를 사용하면 생산성이 증가한다					
5. AI는 전반적으로 나에게 유용하다고 생각한다					

Q9. 다음은 귀하의 인적 특성에 관한 문항입니다.

1. 귀하의 성별	① 남성 ② 여성
2. 귀하의 연령	① 20대 (20~29세) ② 30대 (30~39세) ③ 40대 (40~49세) ④ 50대 (50~59세) ⑤ 60대 이상 (60세~)
3. 귀하의 최종 학력	① 초등학교 졸업 이하 ② 중학교 ③ 고등학교 ④ 대학재학/중퇴(전문대 포함) ⑤ 대학교 졸업 ⑥ 대학원 이상
4. 귀하 가정의 월평균 소득	① 100만원 미만 ② 100~199만원 ③ 200~299만원 ④ 300~399만원 ⑤ 400~499만원 ⑥ 500~599만원 ⑦ 600~699만원 ⑧ 700~799만원 ⑨ 800~899만원 ⑩ 900~999만원 ⑪ 1000만원 이상
5. 귀하의 직업	① 농업, 임업, 어업, 축산업 ② 자영업(종업원 9명 이하의 장사 및 가족 종사자, 개인택시 운전자 등) ③ 판매,영업,서비스직(상점 점원, 세일즈맨 등) ④ 기능,작업직(토목관계의 현장작업, 청소, 수위 등) ⑤ 사무,기술직(일반회사 사무직, 기술직, 초중고교 교사, 항해사 등) ⑥ 경영,관리직(5급 이상의 고위 공무원,기업 부장 이상의 위치, 교장 등) ⑦ 전문직(대학교수, 의사, 변호사, 디자이너 등) ⑧ 자유직(소설가, 화가, 예술인, 프로운동선수 등) ⑨ 전업주부 ⑩ 학생 ⑪ 무직 ⑫ 기타
6. 귀하의 거주지	① 서울특별시 ② 부산광역시 ③ 인천광역시 ④ 대구광역시 ⑤ 대전광역시 ⑥ 광주광역시 ⑦ 울산광역시 ⑧ 경기도 ⑨ 경상북도 ⑩ 경상남도 ⑪ 전라북도 ⑫ 전라남도 ⑬ 충청북도 ⑭ 충청남도 ⑮ 강원도 ⑯ 제주도 ⑰ 세종특별자치시

ABSTRACT

A Study on the Impact of Artificial Intelligence
Technology on Intention to Use
Using Technology Acceptance Model
– Focused on the Mediating Effects of AI Literacy –

Yoon, Jong-Hyun

Major in ESG Management Consulting

Dept. of Knowledge Service & Consulting

Graduate School of Knowledge Service
& Consulting

Hansung University

Since the recent emergence of ChatGPT, a generative AI, Artificial Intelligence technology has been spreading globally in various fields such as businesses, schools, and household, and the rapidly evolving AI technology has brought both positive and negative effects. In the near future society where humans and AI technology coexist, understanding and utilizing AI technology will be an essential skill from the user's perspective. There is a growing need for AI Literacy, defined as the ability to independently and critically evaluate

the information provided by AI and interact with AI technologies. Previous studies have mostly focused on the technological impact of AI, with research on the relationship between AI Literacy and the use of AI technologies being limited.

This study uses the technology acceptance model to understand the technology acceptance process and Intention to Use of users who accept new technologies and services. In addition, this study aims to explain the relation between the level of AI Literacy, represented by the capabilities to understand, utilize, and critically analyze AI technology, and the Intention to Use. Specifically, a structural equation model was built to test the effect of the following variables on the user's intention to continue using AI. Perceived Usefulness and Perceived Ease of Use were set as variables of the technology acceptance model for AI, and the factor of AI Literacy as a mediator variable.

The analysis of this study is based on the online survey results in 2024, December among 312 adults aged 20 and over. The empirical analysis was conducted using the survey data on AI technology. The main findings of the study are as follows.

Among the main variables of the Technology Acceptance Model, Perceived Ease of Use positively affects Perceived Usefulness and Intention to Use, and AI Literacy has a positive effect on Perceived Usefulness. Perceived Usefulness has a positive effect on AI Literacy level and Intention to Use. AI Literacy level has a positive effect on Intention to Use. In addition, this study examines whether AI Literacy level has an indirect effect through the mediating role of AI Literacy level in influencing Perceived Usefulness and Perceived Ease of Use, and determined that only Perceived Usefulness has a significant effect.

This findings have valuable insights in that it attempts the application and extension of the existing theoretical model of AI Literacy, which has become increasingly important among the factors affecting the Intention to Use of AI technology, and empirically verifies the mediating effect of AI Literacy level on users' acceptance of AI technology. Moreover, considering that the rapid development of AI technology is expected in the future, the empirical results of this study can be used as a basis for policy formulation to address the issue of polarization of AI technology users and the AI Literacy gap.

Based on these research results, the following policy implications can be drawn. First, with the advancement of AI technology, new forms of AI Literacy gaps by social class are expected to continue to grow. To address this, it is necessary to strengthen critical AI Literacy skills in a more comprehensive way than solely literacy-related ability to utilize information devices. The concept of diversity, equity, and inclusion (DE&I) is an essential prerequisite for strengthening AI Literacy capabilities and reducing the AI Literacy gap.

Second, it is necessary to spread positive perceptions about the ease of use and Perceived Usefulness of AI in order to expand the use of AI technologies, as shown in the case of cell phones, which are currently used in everyday life by people of all ages without any inconvenience. In particular, two out of three citizens are most concerned about the unemployment and labor insecurity that is expected to worsen due to AI technology. To overcome such negative perceptions, active countermeasures are needed through social consensus among countries, companies, and people, such as agreeing on and building social safety nets and retraining systems for workers.

Third, searching ways to improve the reliability and usability of AI

technology by solving technical issues related to personal information protection and AI technology errors is crucial. To enhance user trust and utilization of AI technology, it is necessary to solve technical issues related to personal information protection and AI errors. This includes measures such as strengthening transparency of AI algorithms so that users can understand how AI works, verifying and evaluating AI learning data to improve AI performance and accuracy. It is also necessary to ensure the rights and safety of users by establishing legal systems and regulations to prevent and address AI technology side effects. These efforts will contribute to improving user trust and utilizing AI technology more effectively.

The limitations of this study are that one, the survey data is restricted to users in Korea, two, the survey was conducted at a specific point in time and may not fully reflect changes in perceptions and behavior patterns over time, and finally, due to the limited sample size of the online survey, the survey questions and data may not reflect all influencing variables related to AI technology. These points suggest that future research directions related to the acceptance and satisfaction of AI use include research on AI Literacy concepts, competencies, and utilization applied to legal, ethical, and social issues.

【Key words】 Artificial Intelligence(AI), AI Literacy, TAM(Technology Acceptance Model), Perceived Ease of Use, Perceived Ease of Use, Intention to Use