

생성형 AI 특성이 과업 성과와 지속 사용 행동에 미치는 영향: SOR과 TTF 이론을 중심으로

The Effects of Generative AI Characteristics on Task Performance and Continued Usage: Based on SOR Framework and TTF Theories

박현선(주저자) · 김상현(교신저자) · 이민영(공저자)

Hyunsun Park(First Author) · Sanghyun Kim(Corresponding Author) · Minyoung Lee(Co-Author)

경북대학교 경영학부 Contract Professor, BK21, School of Business Administration, Kyungpook National University(sunny09@knu.ac.kr)

경북대학교 경영학부 Professor, School of Business Administration, Kyungpook National University(ksh@knu.ac.kr)

경북대학교 경영학부 Ph.D. School of Business Administration, Kyungpook National University(bibianna0910@naver.com)

본 연구는 과업 수행에 생성형 AI 활용이 증가함에 따라 기존 AI와 차별화되는 특성이 사용자 만족과 과업-기술 적합성에 미치는 영향을 살펴보고자 하였다. 또한, 생성형 AI 사용자의 만족과 과업-기술 적합성이 과업성과와 지속 사용에 미치는 영향을 살펴보았다. 이를 위해 생성형 AI의 대표적 서비스인 ChatGPT 사용자를 대상으로 설문조사를 실시하였으며, 총 315부의 설문을 취합하여 AMOS 29.0을 이용해 분석하였다. 연구 결과, ChatGPT의 특성 중 의인화를 제외한 개인화, 지능성, 희소성, 정확성은 사용자 만족과 과업-기술 적합성에 유의한 영향을 미치는 것으로 확인되었다. 또한, 과업-기술 적합성은 사용자 만족, 과업성과, 지속 사용 의도에 모두 유의미한 영향을 미치는 것으로 확인되었으며, 사용자 만족은 과업성과와 지속 사용 의도에 유의미한 영향을 미치는 것으로 확인되었다. 이러한 본 연구의 결과는 생성형 AI 관련 연구의 이론적 범위를 넓히고 과업 수행에 생성형 AI를 도입하려는 조직에 유의미한 시사점을 제공할 수 있을 것이다.

주제어: 생성형 AI, ChatGPT, S-O-R 프레임워크, 과업-기술 적합성, 사용자만족, 과업성과, 지속사용의도

With the increasing use of Generative AI in actual task performance, this study aims to investigate how the characteristics that differentiate Generative AI from traditional AI influence user satisfaction and task-technology fit, based on the Stimulus-Organism-Response framework. Additionally, the study examines the effects of user satisfaction and task-technology fit on task performance and continued usage intention. To analyze the proposed research model, a survey was conducted with users of ChatGPT, a representative Generative AI service. A total of 315 valid responses were collected, and structural equation modeling was performed using AMOS 29.0. The results show that among the characteristics of ChatGPT (Personalization, Intelligence, Rareness, Accuracy) significantly influence both user satisfaction and task-technology fit, while Anthropomorphism has a significant effect only on user satisfaction. Furthermore, task-technology fit significantly affects user satisfaction, task performance, and continued usage intention, and user satisfaction also has a significant impact on task performance and continued usage intention. These findings contribute to expanding the theoretical scope of research on Generative AI and offer practical implications for organizations seeking to adopt generative AI to enhance task performance.

Keyword: Generative AI, ChatGPT, S-O-R framework, Task-Technology Fit, Satisfaction, Task Performance, Continued Usage Intention

최초투고일: 2025. 06. 27

수정일: (1차: 2025. 07. 16)

게재확정일: 2025. 07. 31

Copyright 2025 THE KOREAN ACADEMIC SOCIETY OF BUSINESS ADMINISTRATION

This is an open access article distributed under the terms of the Creative Commons Attribution License 4.0, which permits unrestricted, distribution, and reproduction in any medium, provided the original work is properly cited.

I. 서론

2022년 11월 미국의 인공지능 연구기관인 오픈AI(Open AI)가 대화 기반의 생성형 인공지능(Generative AI) 서비스인 ChatGPT를 공개한 이후, 생성형 AI 기술은 전 세계 경제·기술 흐름에 중심으로 부상했으며, 이를 둘러싼 글로벌 경쟁도 가속화되고 있다. ChatGPT는 출시 약 2년 반 만에 연간 반복 매출(Annual Recurring Revenue, ARR) 100억 달러를 돌파했으며, 2025년 3월 기준 주간 활성 이용자 수(Weekly Active User, WAU)는 5억 명을 넘어서는 등 괄목할 만한 성과를 거두고 있다(김소연, 2025). 특히, 최근에는 GPT-4o 기반의 이미지 생성 기능이 탑재되면서 사용자들 사이에서 '지브리화(Ghiblify)'라는 유행어가 등장하는 등 문화적 파급력도 확산하고 있으며 이 영향으로 한 달 사이 이용자 수는 거의 두 배 가까이 증가한 것으로 나타났다. 국내의 경우 월간 활성 이용자 수(Monthly Active Users, MAU)가 1,000만 명을 초과하는 등 대중적 확산이 매우 빠르게 이루어지고 있다. 오픈 AI는 2024년 약 37억 달러의 매출을 기록한 데 이어 2025년에는 전년 대비 세 배 이상 증가한 127억 달러를, 2026년에는 그보다 두 배가 넘는 294억 달러의 매출을 달성할 것으로 전망되고 있다(선담은, 2025).

이러한 ChatGPT와 같은 생성형 AI는 사용자의 요구에 따라 텍스트, 이미지, 오디오 등 다양한 형태의 결과물을 능동적으로 생성하는 기술로 기존의 딥러닝 기반 AI보다 더 진화된 형태의 기술로 평가받고 있다. ChatGPT는 생성형 AI 기술이 구현되어 상용화된 대표적 사례로 주목받고 있으며, 구글(Google)의 Gemini, 메타(Meta)의 Llama, 마이크로소프트(Microsoft)의 Copilot, 이미지 생성 기반의 DALL-

E2, Midjourney, 오디오 생성 기반의 Suno, Udio 등 다양한 생성형 AI 서비스들이 계속해서 등장하고 있다(양지훈&윤상혁, 2023).

생성형 AI는 기존의 AI 기술보다 더 창의적이면서 사용자에게 맞춤형 결과물을 빠른 속도로 생성한다. 단순한 정보 검색이나 처리에만 국한되지 않고, 사용자와의 상호작용을 통해 맥락을 이해하고 문제 해결을 돕거나 맞춤형 창작물을 생성하고 나아가 의사소통까지도 가능하다. 특히, 인간의 고유한 영역으로 여겨졌던 창작이라는 영역에 AI가 진입해 불과 몇 초 만에 상상하지 못한 결과물을 만들어낸다는 점에서 사회적 파급력이 매우 큰 기술로 인식되고 있다. 또한, 텍스트에만 국한되지 않고 이미지, 오디오, 비디오 등 다양한 형태의 콘텐츠를 생성한다. 이는 인간의 창작 행위 자체에 대한 정의를 재구성하고 콘텐츠 생산과 소비 방식 전반에 구조적인 변화를 불러왔으며, 전 세계 비즈니스 생태계에도 큰 변화를 불러오는 파괴적 혁신을 가져왔다. 이제 생성형 AI는 기술 혁신과 함께 개인뿐만 아니라 조직, 문화 등 디지털 사회에 핵심적인 역할을 할 기술로 평가받고 있다.

생성형 AI가 학습이나 업무 등 실제 과업 수행에 유용한 도구로 활용되는 사례가 많아지면서 생성형 AI 수용 행동과 사용 경험에 대한 이해의 필요성도 점차 커지고 있다. 이에 국내·외에서 생성형 AI 수용과 활용에 관한 연구가 활발히 진행 중이다(선덕길&유재현, 2024; 안승규&안현철, 2024; Albayati, 2024; Biswas et al., 2025; Gupta and Mukjerjee, 2024; Lee et al., 2025; Seifdar and Amiri, 2025). 연구의 관점도 초기에는 기술의 기능적 측면이 주를 이루었으나 현재는 기술 수용과 사용에 관한 연구로 확장되고 있다(Duong et al., 2024; Jo, 2024; Lam, 2025; Zhou and Ma, 2025). 이들 연구는 주로 생성형 AI의 시스템 품질, 정보 품질,

서비스 품질, 사용자 특성, 사회적 영향 등이 신뢰, 만족, 가치 등을 매개로 사용자 행동에 영향을 미친다는 것을 확인하였다(Kang et al., 2024; Lee et al., 2025; Zhou and Ma, 2025).

정보기술의 지속 가능성과 성공을 확보하기 위해서는 사용자의 지속 사용 행동에 영향을 미치는 요인을 파악하는 것이 필수적이다. 특히, 생성형 AI가 빠르게 확산하고 사용자가 급격히 증가하고 있는 시점에 생성형 AI 지속 사용 행동에 대해 살펴볼 필요가 있다. 또한, 생성형 AI가 과업 수행에 활용되는 사례가 많아지고 있다는 점에서 과업 성과 측면에 대해서도 이해할 필요가 있다. 이에 본 연구는 다음의 연구 질문(research question)을 제안하고 이에 대한 답을 찾고자 한다.

RQ1: 다른 서비스와 차별화되는 생성형 AI의 특성은 사용자의 지속 사용 행동에 어떤 영향을 미치는가?

RQ2: 다른 서비스와 차별화되는 생성형 AI 특성은 사용자의 과업 성과에 어떤 영향을 미치는가?

이와 같은 연구 질문에 대한 답을 찾기 위해 먼저 본 연구는 생성형 AI 서비스 중 대중적으로 활발히 사용되고 있는 ChatGPT를 연구 대상으로 설정하고자 한다. 다음으로 정보 검색 플랫폼이나 기존 AI와 차별화되는 특성들을 제안하여 ChatGPT에 대한 사용자 만족과 과업-기술 적합성에 미치는 영향을 살펴보고 더 나아가 지속 사용 행동과 과업 성과에는 어떤 영향을 미치는지를 확인하고자 한다. 특히, 본 연구는 기존 연구들이 단일 이론을 중심으로 생성형 AI 수용을 설명하고 있는 것에서 한 단계 더 나아가, 인간의 행동을 환경 자극과 유기체 그리고 반응의 흐름으로 설명하는 S-O-R(Stimulus-Organism-

Response) 프레임워크와 기술의 특성이 과업의 요구에 적합할 때 기술 수용이 촉진된다는 것을 설명하는 과업-기술 적합(Task-Technology Fit) 이론을 함께 적용하고자 한다. S-O-R 프레임워크는 사용자의 감정, 인지 등 내적 반응과 이를 통해 나타나는 행동을 통해 기술의 수용 과정을 설명할 수 있는 이론적 틀을 제공한다는 점에서 사용자 경험과 반응이 중요한 생성형 AI 사용에 관한 연구에 적합하다. TTF 이론은 기술의 특성과 과업 간의 적합성을 통해 실제 과업 수행에서 기술의 효과성과 효율성을 이해하는데 유용한 이론으로, 최근 생성형 AI가 정보 검색을 넘어 과업 수행에 중요한 도구로써 활용되는 상황을 설명하는데 적절하다. 이에 본 연구는 이 두 이론을 통합적으로 적용하여 ChatGPT를 대상으로 생성형 AI의 특성이 지속 사용 행동과 과업 성과에 어떠한 영향을 미치는지를 살펴보고자 한다. 이러한 통합적 접근은 생성형 AI의 사용 맥락을 보다 입체적으로 이해할 수 있는 이론적 틀을 제공함과 동시에 관련 연구의 이론적 범위를 확장하고 실무적인 시사점 또한 제공할 수 있을 것으로 기대한다.

II. 이론적 배경

2.1 생성형 AI와 ChatGPT

생성형 AI는 기계학습과 딥러닝을 기반으로 대규모 데이터에 대한 학습을 지속하면서 사용자의 지시에 따라 사용자가 원하는 새롭고 독창적인 콘텐츠(텍스트, 이미지, 오디오, 비디오 등)를 창작해 내는 인공지능 기술 또는 서비스를 말한다(선덕길&유재현, 2024; Gupta et al., 2024). 대형 언어모델(Large

Language Model, LLM)이나 이미지 생성 모델(Image Model, IGM)을 활용해 새로운 창작물을 생성하는 모든 기술을 포함한다. 기존의 딥러닝 기반의 AI는 주로 학습 데이터를 분석하고 분류(classification) 및 예측(prediction) 등의 기능 수행에 초점을 맞추었다. 반면, 생성형 AI는 사용자가 입력한 데이터를 바탕으로 새로운 형태의 콘텐츠나 해결책을 생성하는 데 중점을 둔다는 점에서 기존 AI와 뚜렷한 차이를 보인다.

김수경(2024)은 생성형 AI가 방대한 데이터로부터 학습을 통해 텍스트, 이미지, 디자인, 코드 등의 새로운 정보를 생성할 수 있으며, 기존 AI와 비교했을 때 대화형 상호작용, 콘텐츠 생성 능력, 맞춤형 서비스 제공, 적응성과 확장성, 자기 학습 능력 등의 특성이 있다고 설명한다. 즉, 생성형 AI는 인간과의 상호작용을 통해 실시간으로 복잡한 작업을 수행할 수 있으며, 창의적이고 고품질의 콘텐츠를 생성함으로써 새로운 가치를 창출한다. 또한, 사용자 개개인의 선호와 행동을 분석하여 맞춤화되고 최적화된 콘텐츠를 제공하고 주어진 데이터나 상황에 따라 자동으로 해결책이나 콘텐츠를 조정함으로써 변화에 빠르게 대응할 수 있다. 나아가, 새로운 데이터를 지속적으로 학습함으로써 시간이 지날수록 더욱 정확하고 효과적인 결과를 도출할 수 있다는 점도 주요 특징 중 하나이다.

생성형 AI는 2014년 적대적 생성 신경망(Generative Adversarial Networks, GAN)의 등장 이후 지속적으로 발전해 왔으며, 오픈AI(Open AI)가 GPT-3.5 기반 언어모델을 활용해 발표한 ChatGPT의 등장으로 본격적인 주목을 받기 시작했다. 특히, 2022년 출시된 ChatGPT는 LLM을 기반으로 한 대화형 생성형 AI 챗봇으로, 생성형 AI 기술이 상용화된 대표적인 사례이자 가장 강력한 생성 모델 중 하나로 평

가된다(Gupta et al., 2024). LLM은 대규모 언어 데이터를 학습하여 문맥을 이해하고, 이에 기반한 적절한 단어와 문장을 예측, 생성함으로써 자연스러운 텍스트를 생성하는 모델로, 챗봇을 포함한 다양한 텍스트 생성 서비스에 널리 활용되고 있다. ChatGPT는 문장의 맥락을 파악하여 단어의 문법적, 의미적 특성을 이해하고 그에 맞는 텍스트를 생성함으로써 사용자와 자유로운 대화가 가능하다. 또한, 인터넷의 방대한 정보와 공개적으로 접근이 가능한 데이터를 학습하여 사용자가 요구하는 다양한 질문에 답변할 수 있는 능력을 갖추고 있다. 이런 점에서 ChatGPT는 다양한 작업을 수행할 수 있는 능력을 보유하고 있어 다양한 분야에서 활용이 가능하다는 이점을 가진다.

생성형 AI가 급속히 확산하면서 글로벌 빅테크 기업을 비롯한 다양한 기업들이 관련 시장에 진입하여 자체 플랫폼과 서비스를 개발 및 출시하고 있다. 대표적으로 구글(Google)의 Gemini(구 Bard), 메타(Meta)의 Llama, 마이크로소프트(Microsoft)의 Bing Copilot 등이 있으며, 국내에서는 네이버(Naver)가 한국어에 특화된 OCEAN 언어모델을 기반으로 CLOVA X를 개발하여 상용화한 바 있다. 또한, ChatGPT와 같은 텍스트를 기반으로 하는 생성형 AI뿐만 아니라 이미지, 오디오, 비디오 등을 생성하는 다양한 형태의 생성형 AI가 등장하면서 공학, 의료, 제조, 금융, 건축 등 여러 분야에서 광범위한 활용 가능성을 지닌 핵심 기술로 주목받고 있다. 이미지 기반 생성형 AI는 주어진 데이터의 잠재 표현을 학습하여 기존 이미지를 개선하거나 변형하고 고품질의 정확한 이미지를 생성하는 AI를 말한다. 대표적인 예로 DALL-E2, Midjourney, Adobe Firefly, Playground AI 등이 있다. 오디오 기반 생성형 AI는 다양한 소리와 음악을 만들어내거나 자연스럽게 음성을 합성할 때 활용되며 Suno, Udio, AIVA 등이

있다. 비디오 기반 생성형 AI는 영상 콘텐츠를 편집하거나 새로운 영상을 생성하는 것으로 Sora, Pika, Pictory 등이 대표적이다.

이처럼 생성형 AI는 다양한 서비스의 출시와 함께 활용 범위가 빠르게 확대되면서 여러 산업 전반으로 적용이 확대되고 있으며, 나아가 개인이나 기업의 업무 효율성 향상과 조직 운영의 혁신에 기여할 유용한 기술로 인식되고 있다(Seifdar and Amiri, 2025). 이에 따라 제조, 의료, 교육, 국방, 금융, 예술 등 다양한 분야에서 생성형 AI에 대한 연구가 활발히 진행되고 있으며(이한신&김관수, 2019; Albayati, 2024; Lee et al., 2025; Tiwari et al., 2023; Verma et al., 2025), 경영학 분야에서도 관련 연구가 지속적으로 확산되고 있다. 경영학 연구는 주로 생성형 AI를 활용하는 개인의 행동에 영향을 미치는 요인, 그리고 조직, 업무 성과, 산업별 환경에 미치는 영향에 주목하고 있다(선덕길&유재현, 2024; 이원준, 2024; Biswas et al., 2025; Duivenvoorde, 2025; Duong et al., 2024; Gupta and Mukjerjee, 2024; Jo, 2024). 예를 들어, Biswas et al.(2025)은 학생들 ChatGPT의 사용 경험과 추천 행동에 영향을 미치는 요인을 분석하였으며, 이를 통해 학문적 관점에서 생성형 AI 활용에 관한 보다 심층적인 연구의 필요성을 제기하였다. 또한, Duong et al.(2024)은 생성형 AI가 특히 교육 분야에서 큰 영향을 미치고 있음을 강조하면서, ChatGPT의 지속 사용 행동에 영향을 미치는 요인을 실증적으로 규명하였다.

Gupta and Mukjerjee(2024)는 생성형 AI가 소비자의 정보 탐색 방식에 변화를 초래할 것으로 전망하며, 생성형 AI 플랫폼 채택에 영향을 미치는 요인을 분석한 연구를 수행하였다. 이들은 생성형 AI의 활용이 정보 과부하를 완화할 수 있다고 보았다. Jo(2024)는 생성형 AI 유료 버전에 대한 구독 의도

를 살펴보는 연구에서 시스템 품질, 서비스 품질, 지각된 지능화가 만족도와 목표 일치성을 매개로 하여 구독 행동에 유의미한 영향을 미친다는 점을 실증적으로 확인하였다.

Seifdar and Amiri(2025)는 생성형 AI의 발전이 조직 구조에 변화를 일으키고, 이를 통해 새로운 기회를 창출할 수 있다고 주장하였다. 특히, 생성형 AI 도입은 부서 간 협업을 촉진하고 장기적으로 조직에 긍정적인 효과를 가져올 것이라고 보았다. 선덕길&유재현(2024)은 가치기반수용모델(Value-based Adoption, VAM)을 바탕으로 생성형 AI 사용자의 지속 사용 의도에 영향을 미치는 요인을 분석하였으며, 지각된 유용성과 상호작용성이 생성형 AI의 지속적 사용에 긍정적인 영향을 미칠 수 있음을 확인하였다.

이처럼 생성형 AI의 활용 범위가 확대됨에 따라, 사용자 행동에 영향을 미치는 요인을 규명하려는 연구가 국내외에서 활발히 진행되고 있다. 그러나 생성형 AI만의 고유한 특성과 이에 대한 사용자의 내적 상태 변화, 그리고 지속적인 이용 행동 간의 관계를 실증적으로 분석한 연구는 아직 부족한 실정이다. 이에 본 연구는 기존 AI와 구별되는 생성형 AI의 특성이 사용자에게 어떤 내적 반응을 유도하고, 이는 다시 어떠한 행동적 반응으로 이어지는지를 실증적으로 규명하고자 한다.

2.2 S-O-R 이론

S-O-R 프레임워크는 개인의 감정이나 태도, 행동이 외부의 자극에 의해 변화한다는 관점에 기반하여, 환경이 인간 행동에 미치는 영향을 설명하는데 유용한 이론적 모델로 널리 활용되고 있다. Woodworth(1929)는 고전적인 자극-반응 이론(Stimulus-Response Theory)이 인간 행동을 설명하는 데 있

어 내적 변화를 간과한다는 점에서 불완전하고 과학적이지 않다고 지적하였다. 그는 이를 보완하기 위해 자극과 반응 사이에 존재하는 인지적, 정서적 상태를 강조하며 해당 개념을 처음으로 제시하였다. 이후 Mehrabian and Russell(1974)은 환경에 대한 감정적 반응에 초점을 맞춰 이 모델을 확장하였고, Belk(1975)는 기존의 자극-반응 이론에 유기체(Organism) 요소를 추가함으로써 오늘날의 자극(Stimulus)-유기체(Organism)-반응(Response)으로 구성되는 S-O-R 프레임워크를 정립하였다.

S-O-R 프레임워크에서 자극(Stimulus)은 개인의 내적 상태에 변화를 일으키는 외부 환경적 요인을 의미하며, 인간에게 영향을 줄 수 있는 모든 외부적 요소가 자극에 포함될 수 있다. 기존 연구들(Abumalloh et al., 2025; Cheng et al., 2022; Jacoby, 2002)은 이와 같은 자극의 범위를 물리적 환경 요소뿐만 아니라, 특정 대상과의 상호작용에서 인지되는 정보 품질, 시스템 품질, 서비스 품질 등 지각된 기술적 속성까지 확장하여 해석하고 있다.

유기체(Organism)는 인간의 내적 상태를 의미하며, 외부 자극과 행동(반응) 사이에 개입하는 중개 과정을 의미한다. 이는 외부 환경에 영향을 받은 개인의 정신적 반응으로, 감정 상태로 개념화하여 자극은 감정 상태에 영향을 미치고 이를 매개로 최종적인 행동 반응으로 이어진다고 본다. 유기체에는 인지적, 감정적, 생리적 상태뿐만 아니라 지각적 과정까지 포함될 수 있다. 기존의 S-O-R 프레임워크 기반 연구들은 만족, 신뢰, 인지 및 정서적 반응, 몰입, 즐거움, 경험, 지각된 용이성 및 유용성 등 다양한 요인들을 유기체 요소로 제안하고 있으며 이는 사용자의 내적 변화가 행동에 어떻게 영향을 미치는지를 설명하는 데 중요한 변수로 작용한다(Cheng et al., 2022; Premathilake et al., 2025; Uddin et

al., 2025).

마지막으로 반응(Response)은 외부 자극과 유기체 상태에 따른 인간의 자연스러운 행동 반응을 의미하는 것으로 태도적 또는 행동적 반응을 말한다(Premathilake et al., 2025; Jacoby, 2002). 즉, 자극 때문에 유도된 내적 상태의 변화에 따라 나타나는 반응으로서 이용의도, 지속적 이용의도, 구매 행동, 전환행동, 충성도 등의 형태로 나타난다.

S-O-R 프레임워크에서 자극이 처리되는 과정은 감정과 의사결정에 영향을 미치며, 이러한 일련의 절차를 통해 외부 환경 자극에 대한 인간 행동이 구조적으로 설명될 수 있다(Mehrabian and Russell, 1974; Uddin et al., 2025). 또한, 자극이 처리되는 과정은 의식적일 수도 있고 무의식적일 수도 있다. 이처럼 S-O-R 프레임워크는 환경심리학을 대표하는 이론 중 하나로, 기존 행동심리학의 고전적인 자극-반응 모델에서 벗어나 자극을 해석하고 처리하는 인지적, 정서적 요소를 포함함으로써 인간 행동의 복잡성을 이해할 수 있는 체계적인 분석 틀을 제공한다. 이러한 점에서 S-O-R 프레임워크는 인간 행동의 원인을 파악하고 관련 문제를 해결하는데 유용한 이론으로 다양한 분야에서 폭넓게 활용되고 있다.

경영학과 정보시스템(IS) 분야에서도 S-O-R 프레임워크는 기술 사용자의 태도나 행동 반응에 영향을 미치는 요인을 설명하는 데 효과적인 이론으로 활용되고 있다. 기존 연구들은 정보기술의 수용 및 채택, 신뢰, 만족 등과 같은 사용자 태도에 외부 환경 요인이 어떠한 영향을 미치는지를 설명하기 위해 S-O-R 프레임워크를 채택하였다(Cheng et al., 2022; Lee and Chen, 2022; Premathilake et al., 2025; Truong and Chen, 2025; Uddin et al., 2025; Vafaeio-Zadeh et al., 2024). 예를 들어, Lee and Chen(2022)은 모바일 뱅킹 앱 수

용에 관한 연구에서 지능성과 의인화 인식이 자극으로 작용하여 긍정적인 수용 행동을 유도함을 밝혔다. Cheng et al.(2022)은 챗봇 사용 연구에서 사용자가 지각하는 챗봇의 특성이 자극이 되어 신뢰를 형성하는 경로를 설명하였다.

Premathilake et al(2025)은 휴머노이드 소셜 로봇에 대한 사용자 만족 연구에서, 의인화된 특성에 대한 인식이 자극이 되어 가치 지각과 만족이라는 유기체 반응을 통해 최종 행동으로 이어질 수 있음을 S-O-R 프레임워크를 통해 입증하였다. 또한, Uddin et al.(2025)은 개인적 및 환경적 자극이 내적 상태의 변화를 유발하여 디지털 결제 플랫폼의 채택 행동에 영향을 줄 수 있음을 실증적으로 규명하였고, Truong and Chen(2025)은 챗봇의 의인화와 반응성이 사회적 실재감과 공감을 강화하고, 이를 통해 사용자 경험과 지속 이용 의도에 긍정적인 영향을 미친다는 점을 S-O-R 프레임워크를 통해 확인하였다. Vafaei-Zadeh et al.(2024)은 AI 고객 서비스 도입에 영향을 미치는 요인을 분석하기 위해 S-O-R 프레임워크를 활용하였다. 이처럼, S-O-R 프레임워크는 외부 자극이 개인의 내적 상태를 매개로 하여 태도나 행동 반응으로 이어지는 과정을 설명하는 데 유용하며, 경영학과 정보시스템 분야를 포함한 다양한 연구에서 활발히 활용되고 있다.

이처럼 기존 생성형 AI 관련 연구에서도 S-O-R을 적용한 사례가 있다. 그러나 본 연구는 생성형 AI의 특성이 사용자 만족과 과업-기술 적합성을 매개로 하여 지속 사용 행동과 과업 성과에 미치는 영향 구조를 분석하고자 한다. 이러한 변수 관계를 설명하고, 사용자의 행동 반응을 설명하는 데도 S-O-R 프레임워크가 이론적으로 적절하다고 판단하였다. 이에 본 연구는 이를 연구모형 설계에 있어 이론적 기반으로 채택하여 연구를 진행하고자 한다.

2.3 과업-기술 적합성

정보기술의 '사용(utilization)'이나 기술 '적합(fit)'에 관한 기존 연구들은 개인이 업무를 수행할 때 필요로 하는 기능을 정보기술이 얼마나 효과적으로 지원하고, 성과 향상에 기여하는지를 중점적으로 다루어왔다. 이러한 연구들은 주로 두 가지 관점에서 접근되었는데, 첫 번째는 정보기술 사용의 관점으로 사용자의 태도와 신념과 같은 심리적 요인이 실제 사용 행동과 성과에 어떤 영향을 미치는지를 분석하는 것이다. 두 번째는 업무와 기술 간의 적합성 관점으로 정보기술이 업무 요구를 얼마나 잘 충족시키는지 그리고 그 적합성이 업무 성과나 사용 행동에 어떤 영향을 미치는지를 중점적으로 살펴본다.

Goodhue and Thompson(1995)은 정보기술 사용의 관점을 기반으로 한 기존 연구들이 기술 특성과 과업 특성 간의 적합성을 충분히 고려하지 않아 정보기술의 효율적인 활용과 이해에 한계가 있다고 지적하면서 과업-기술 적합성(Task-Technology Fit, TTF) 모델을 제안하였다. TTF는 개인이 과업을 수행할 때 필요한 요구를 정보기술의 기능이 얼마나 잘 지원하는지를 설명하는 이론적 모델로 정의된다. TTF 모델에서 '과업(task)'은 개인이 수행하는 구체적인 업무나 활동을 의미하며, '기술(technology)'은 이러한 업무 수행을 지원하는 기능적 수단을 의미한다. 그리고 이 두 요소와 사용자 간의 상호작용을 설명하기 위해 '적합(fit)'이라는 개념이 포함되며, 이는 기술이 과업의 요구와 어느 정도 일치하는지를 나타낸다(Howard and Rose, 2019; Jeyaraj, 2022). TTF 모델은 과업과 기술 간의 적합성이 기술의 실제 사용과 사용자의 성과 향상에 있어 중요한 매개 역할을 한다고 전제한다. 즉, 기술이 과업 수행에 효과적으로 부합할수록 사용자는 그 기술을 더욱 적극적으로

로 활용하게 되며 이는 결과적으로 향상된 성과와 긍정적인 기술 사용 행동으로 이어질 수 있다(Goodhue and Thompson, 1995; Lam, 2025). 또한, 과업과 기술 간의 간극이 클수록 사용자는 과업과 기술 간의 적합성을 낮게 인지하며 이에 따라 기술에 대한 가치와 중요성에 대한 인식도 감소하게 된다. 따라서 기술의 지속적인 사용은 사용자가 수행해야 할 과업과 기술의 특성 간에 명확한 일치가 존재해야만 가능하다.

TTF 모델을 활용한 연구들은 과업과 기술 간의 적합성이 사용자의 태도와 업무 성과 그리고 사용에 긍정적인 영향을 미친다는 사실을 확인하였다(Jeyaraj, 2022; Wang et al., 2020; Wu and Chen, 2017; Zhang et al., 2025). 예를 들어, Wang et al. (2020)은 TTF를 활용한 헬스케어 웨어러블 기기 사용에 관한 연구에서 기술 특성과 과업 특성 간의 적합성이 높을수록 사용자들의 기술 사용에 긍정적인 영향을 미친다는 점을 확인하였다. 또한, Wu and Chen(2017)은 MOOC 환경에서 과업-기술 적합성이 사용자의 지속적 사용 행동에 유의미한 영향을 미친다는 사실을 확인하였다. Zhang et al.(2025)은 기업용 소셜미디어를 대상으로 한 연구에서 과업 특성과 기술 특성 간의 적합성이 개인의 디지털 성과에 긍정적인 영향을 미친다는 것을 실증분석을 통해 확인하였다.

최근 연구에서는 TTF 모델을 전통적인 정보기술 뿐만 아니라, 최신 기술이나 서비스에도 확장하여 적용하고 있다. 예를 들어, AI 기반 쇼핑 플랫폼이나 AI 고객센터 도입이나 사용의 맥락에서 AI 기술의 의사소통 능력, 지능성, 의인화, 유연성 등이 쇼핑이나 고객센터 과업을 얼마나 잘 지원하는지가 중요한 평가 기준이 된다. 또한, 기술적 특성과 과업 특성이 적합할 경우, 사용자는 더 효율적으로 목표를 달성할 수 있고 사용 행동을 유도하게 된다(Chakraborty

et al., 2025; Vafaei-Zadeh et al., 2024).

ChatGPT에 관한 연구에서도 TTF는 개인의 행동을 설명하는데 유용한 이론적 틀로 활용되고 있다. ChatGPT와 같은 생성형 AI는 강력한 자연어 처리 기술을 활용하여 일상이나 직장 생활 속에서 나타나는 다양한 언어적 행동을 모방하여 사용자의 과업 수행을 효과적으로 지원한다(이원준, 2024). 업무에 생성형 AI를 활용하는 사용자가 증가함에 따라 생성형 AI가 업무에 얼마나 적합한지에 대한 인식의 정도는 기술 수용이나 사용 행동에 영향을 미칠 수 있다(안승규&안현철, 2024; Lam, 2025).

이처럼 일부 선행연구에서 생성형 AI 사용 맥락에서 TTF를 적용한 사례가 있지만, 연구의 대부분은 TTF 변수와 사용행동 간의 관계에만 초점을 맞추고 있고 지속 사용 행동과 과업 성과까지 이어지는 구조를 다룬 연구는 제한적이다. 따라서 본 연구는 TTF를 기존보다 확장된 구조로 재적용하여 ChatGPT와 같은 생성형 AI의 특성이 사용자의 과업 수행에 있어 과업 성과와 지속 사용 행동에 미치는 영향을 함께 분석하고자 한다. 이러한 접근은 생성형 AI의 특성이 실제 업무 수행에 어떻게 기여하는지에 대한 이론 및 실증적 이해를 도모한다는 점에서 의의가 있을 것으로 기대한다.

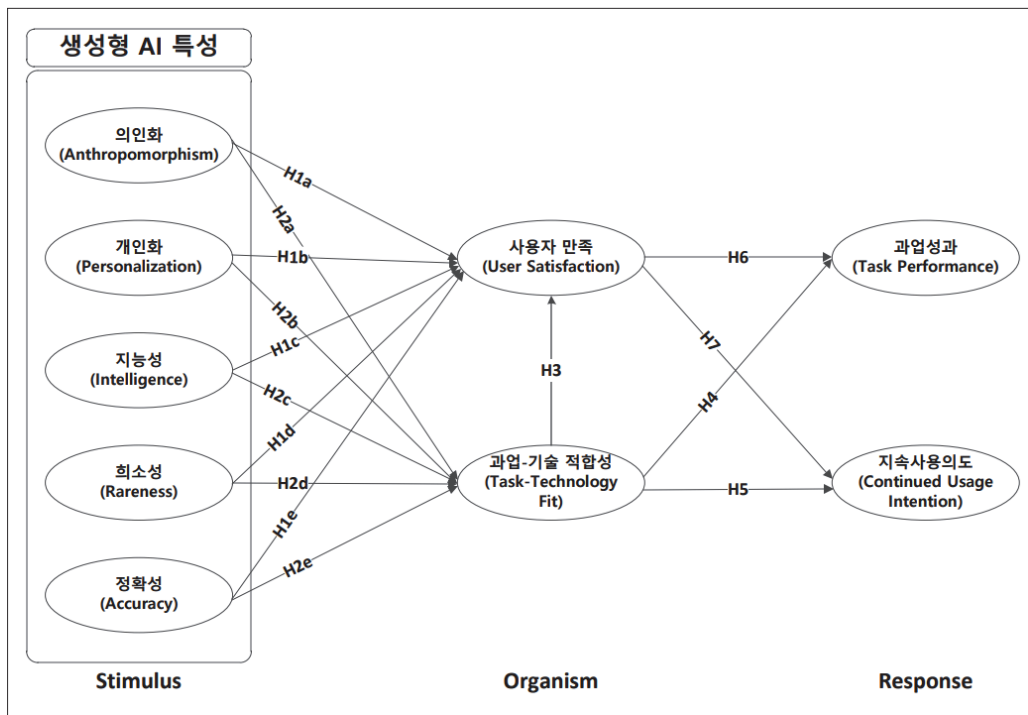
III. 연구모형 및 가설

3.1 연구모형

본 연구의 목적은 ChatGPT와 같은 생성형 AI의 특성이 과업 성과 및 지속 사용 의도에 미치는 영향을 실증적으로 규명하는 데 있다. 이를 위해, 인간의

행동을 이해하는데 유용한 이론적 틀로 알려진 S-O-R 프레임워크를 기반으로 연구모형을 제안하였다. 먼저 외부 자극에 해당하는 변수로는 생성형 AI의 핵심 특성을 도출하였으며, 선행연구에 대한 문헌 고찰을 통해 의인화, 개인화, 지능성, 희소성, 정확성의 다섯 가지 요인을 제시하였다. 다음으로 유기체에 해당하는 내적 심리 및 인지적 상태를 설명하기 위해서는 사용자 만족과 과업-기술 적합성을 변수로 제안하였으며, 이들이 생성형 AI의 특성과 어떤 관계가 있는지를 분석하고자 하였다. 마지막으로 행위적 반응에 해당하는 변수로는 과업 성과와 지속 사용 의도를 제시하였으며, 사용자 만족과 과업-기술 적합성이 이러한 행동적 결과 변수에 미치는 영향을 검증하고자 하였다. 이러한 연구모형 설계를 위해 본 연구는 S-O-R

프레임워크와 TTF 이론을 통합적으로 적용하였다. S-O-R은 외부 자극이 내적 반응인 유기체를 거쳐 행동으로 이어지는 과정을 구조적으로 설명하는데 이점이 있다. 반면, TTF는 기술이 과업을 수행하는데 얼마나 적합한지에 대한 사용자의 인식이 실제 사용과 성과에 미치는 영향을 설명하는데 이점이 있다. 이에 본 연구는 S-O-R 프레임워크를 통해 생성형 AI 특성이 사용자 만족과 같은 내적 반응 그리고 지속 사용 행동으로 이어지는 구조를 설명하고, TTF 이론을 통해 과업-기술 적합성을 함께 살펴봄으로써 생성형 AI 사용 행동과 생성형 AI의 영향을 보다 포괄적으로 이해하고자 하였다. 이러한 논의를 바탕으로 본 연구는 <그림 1>과 같은 연구모형을 제시하고, 이에 따른 연구가설을 도출하였다.



〈그림 1〉 연구모형

3.2 가설설정

3.2.1 생성형 AI 특성 요인

ChatGPT와 같은 생성형 AI와 같은 기술의 핵심은 사용자와의 자유로운 상호작용을 통해 요구사항을 이해하고 이를 분석, 정리하여 적절한 결과물을 생성하는 데 있다(Kang et al., 2024). 기존의 AI 기술은 주로 단순한 정보 검색이나 정형화된 질문에 대한 응답에 초점을 맞추었으며, 질문의 맥락을 파악하거나 복잡한 요청에 유연하게 대응하는 데에는 한계가 있었다(Gupta et al., 2024). 반면, 생성형 AI는 마치 사람과 대화하듯 자연스럽게 맥락을 이해하고, 단순한 정보 검색으로는 얻기 어려운 통찰력 있는 답변이나 창의적인 콘텐츠를 능동적으로 제공할 수 있으므로 차별화된 사용자 경험을 가능하게 한다. 이에 본 연구는 생성형 AI에 대한 이론적 문헌 고찰을 바탕으로 기존의 AI와 차별화되는 생성형 AI의 고유한 특성을 의인화, 개인화, 지능성, 회소성, 정확성의 다섯 가지 변수로 제안하고 이를 S-O-R 프레임워크에서 외부 자극에 해당하는 요인으로 보았다. 이는 S-O-R 프레임워크에 관한 선행연구에서 외부 자극의 범위를 특정 대상과의 상호작용에서 인지되는 요소로 정보 품질, 시스템 품질, 서비스 품질 등 지각된 기술적 속성까지 확장하여 해석하고 있기 때문이다(Cheng et al., 2022; Jacoby, 2002).

먼저 의인화(Anthropomorphism)란 사용자가 ChatGPT에 대해 인간과 유사한 행동양식, 감정, 동기 등이 부여되어 있다고 인식하는 정도를 의미하며, 이는 곧 사용자가 해당 기술을 얼마나 인간처럼 인식하는지를 반영한다(Gupta and Mukherjee, 2024; Moussawi et al., 2021; Zhou and Ma, 2025). 한편, 지능성(Intelligence)은 사용자가

ChatGPT를 사용할 때 이를 똑똑하고 유능하다고 평가하는 주관적인 인식을 의미하며 ChatGPT의 지적 역량에 대한 사용자의 인지적 판단을 나타낸다. 생성형 AI와 서비스 로봇에 대한 선행연구들은 이러한 기술이 의인화와 지능성이라는 두 가지 핵심 특성을 보인다고 보고하고 있다. AI 기반 시스템은 종종 인간의 감정이나 행동을 모방하도록 설계되며 이러한 특성은 의인화로 정의된다. 동시에, 자율적이고 효율적으로 기능을 수행하며 사용자 요청에 따라 적절한 결과를 생성할 수 있는 능력은 지능성으로 간주한다. 사용자는 생성형 AI가 인간과 유사한 행동을 보이거나 자연스러운 대화를 제공할 때 심리적 안정감을 더 크게 느끼며 기술에 대한 수용 태도 또한 긍정적으로 변화하는 경향이 있는 것으로 나타났다(Lee and Chen, 2022). 더불어 사용자가 생성형 AI를 인간과 유사한 존재로 인식할수록 해당 기술에 대한 만족도나 과업-기술 적합성 역시 보다 긍정적으로 평가하는 경향이 있다(Lee and Chen, 2022; Zhou and Ma, 2025). 또한, 생성형 AI가 사용자의 요청을 자율적으로 분석하고 이에 적절하게 반응할 수 있을 때, 사용자는 해당 기술이 높은 수준의 지능을 갖추고 있다고 인식하며 이를 통해 문제를 해결하거나 정보를 얻고자 하는 기술 활용 의지 또한 증가하는 것으로 보고되고 있다(Moussawi et al., 2021). 안승규&안현철(2024)은 생성형 AI가 마치 인간과 상호 작용하는 것과 같은 경험을 제공하는 특성과 사용자와의 대화를 기억하고 메시지와 맥락을 파악하여 유연한 대화를 전개하는 특성이 과업-기술 적합성에 유의미한 영향을 미침을 확인하였으며, Lee and Chen(2022)은 의인화와 지능화가 과업-기술 적합성에 유의미한 영향을 미침을 확인한 바 있다.

개인화(Personalization)는 ChatGPT가 제공하는 정보와 콘텐츠가 사용자 개인의 선호와 필요에

맞춰져 있다고 사용자가 지각하는 정도를 의미하며, 이는 곧 사용자가 자신에게 맞춤형 응답을 받고 있다고 느끼는 인식을 반영한다(이진&오현정, 2024; Gupta, 2025; Zhou and Ma, 2025). 생성형 AI의 핵심 기능은 사용자의 구체적인 요구에 따라 맞춤형 콘텐츠를 생성하는데 있으며, 이는 기술 활용의 본질적 목적 중 하나로 간주된다. Zhou and Ma (2025)는 사용자가 자신의 다양한 요청에 대해 적절하고 개인화된 정보를 제공받고 있다고 인식할수록 그에 따른 만족감 또한 높아짐을 실증분석을 통해 확인하였으며, Chung et al.(2020)은 챗봇 서비스의 개인화 요소가 만족을 증가시킴을 확인하였다. 이처럼 개인화는 생성형 AI를 사용하려는 주요 동기 요인으로 작용하며, 기술 수용 및 지속 사용 의도에 긍정적인 영향을 미치는 것으로 보고되고 있다(이진&오현정, 2024; Zhou et al., 2022).

희소성(Rareness)은 ChatGPT가 제공하는 정보나 콘텐츠가 다른 AI 시스템이나 일반적인 정보 플랫폼에서는 쉽게 얻을 수 없는 독창성과 창의성을 지니고 있다고 사용자가 인식하는 정도를 의미한다(안승규&안현철, 2024; Zhou et al., 2022). 이는 곧 생성형 AI가 생성하는 콘텐츠가 얼마나 희귀하고 창의적으로 인식되는지를 나타내며, 다른 정보원과 생성형 AI의 차별성을 부각시키는 핵심 요인으로 작용한다. 사용자는 생성형 AI가 제공하는 정보나 콘텐츠에 희소성이 있고 참신하다고 인식할수록 이를 더욱 긍정적으로 평가하며, 이는 만족도와 지속 사용 의도에 긍정적인 영향을 미칠 수 있다. Choi et al. (2017)은 개인화 추천시스템에 관한 연구에서 기술이 주는 콘텐츠가 참신하고 희소하다고 느낄수록 만족감이 높아짐을 확인하였다.

마지막으로 정확성(Accuracy)은 ChatGPT가 제공하는 정보와 콘텐츠가 사용자의 요구와 일치하며

정확한 정보로 인식되는 정도를 의미한다(Gupta, 2025; Zhou and Ma, 2025). 이는 곧 정보의 일관성 및 신뢰성과 밀접하게 관련되며, 생성형 AI가 제공하는 콘텐츠가 실제로 정확하고 신뢰할 수 있으며, 사용자의 요구를 충족시키고 문제 해결에 실질적인 도움을 주는지를 나타낸다(Gupta, 2025). 구체적이고 정확한 콘텐츠는 과업 수행에 직접적인 기여를 함으로써 사용자 효율성을 높이고 만족도를 증대시킬 수 있으며, 나아가 생성형 AI의 성능에 대한 신뢰감을 강화하는데 기여한다. 이러한 정확성에 대한 인식은 해당 기술을 긍정적으로 평가하고 지속적으로 사용하는 주요 동기 요인으로도 작용할 수 있다(Gupta, 2025). Zhou and Ma(2025)는 생성형 AI의 정확성이 만족과 문제 해결에 긍정적인 영향을 미침을 실증분석을 통해 확인하였으며, 박기호&이군호(2024)는 ChatGPT의 정확성과 같은 정보품질이 만족을 부분적으로 증가시켜줄 수 있음을 확인하였다.

본 연구는 선행연구에 대한 이론적 고찰을 바탕으로 제안된 생성형 AI의 다섯 가지 특성이 사용자 만족(User Satisfaction)과 과업-기술 적합성(TTF)에 긍정적인 영향을 미칠 것으로 판단하였다. 이에 따라, 본 연구는 각 변수 간의 관계를 실증적으로 검증하기 위해 다음과 같은 연구가설(H1a~H2e)을 설정하였다.

H1a: 의인화는 사용자 만족에 정(+)의 영향을 미칠 것이다.

H1b: 개인화는 사용자 만족에 정(+)의 영향을 미칠 것이다.

H1c: 지능성은 사용자 만족에 정(+)의 영향을 미칠 것이다.

H1d: 희소성은 사용자 만족에 정(+)의 영향을 미칠 것이다.

H1e: 정확성은 사용자 만족에 정(+)의 영향을 미칠 것이다.

H2a: 의인화는 과업-기술 적합성에 정(+)의 영향을 미칠 것이다.

H2b: 개인화는 과업-기술 적합성에 정(+)의 영향을 미칠 것이다.

H2c: 지능성은 과업-기술 적합성에 정(+)의 영향을 미칠 것이다.

H2d: 회소성은 과업-기술 적합성에 정(+)의 영향을 미칠 것이다.

H2e: 정확성은 과업-기술 적합성에 정(+)의 영향을 미칠 것이다.

3.2.2 과업-기술 적합성

과업-기술 적합성은 특정 기술이 과업의 요구사항과 어느 정도 일치하는지를 나타내며, 사용자가 해당 기술이 과업 수행에 적합하다고 인식할수록 그 기술에 대한 가치와 중요성에 대한 인식 또한 높아지는 것으로 알려져 있다(Goodhue and Thompson, 1995; Lam, 2025). 기술 수용에 대한 긍정적인 태도와 지속적인 사용 의도 역시 기술과 과업 간의 특성적 일치가 전제되어야 가능한 것으로 적합성이 높을수록 사용자는 해당 기술을 더욱 적극적으로 수용하고 활용하는 경향을 보인다(Jeyaraj, 2022; Wu and Chen, 2017).

ChatGPT와 같은 생성형 AI 기술의 경우에도 의인화, 유연성, 지능성, 개인화 등의 특성이 사용자의 과업 수행을 얼마나 효과적으로 지원하고 보완하는지가 중요한 평가 기준이 된다. 사용자가 해당 기술이 자신의 과업을 원활히 수행하도록 실질적인 도움을 제공한다고 인식할수록 이는 사용자 만족도와 지

속적 사용 행동을 유도할 뿐만 아니라 궁극적으로 과업 성과에도 긍정적인 영향을 미치는 것으로 나타났다(Chakraborty et al., 2025; Vafaei-Zadeh et al., 2024; Zhang et al., 2025). 여기서 과업 성과란 사용자가 자신의 과업을 수행하는 과정에서 나타나는 결과로, 과업에 적합한 생성형 AI를 활용할 경우 보다 효과적이고 긍정적인 성과를 도출할 수 있음을 의미한다. 이와 같은 선행연구에 대한 이론적 고찰을 통해 본 연구는 사용자가 인지하는 생성형 AI의 과업-기술 적합성이 사용자 만족, 과업 성과 그리고 지속 사용 의도에 긍정적인 영향을 미칠 것으로 판단하였다. 이러한 가정에 따라 변수 간의 관계를 실증적으로 검증하기 위해 다음과 같은 연구가설(H3~H5)을 설정하였다.

H3: 과업-기술 적합성은 사용자 만족에 정(+)의 영향을 미칠 것이다.

H4: 과업-기술 적합성은 과업성과에 정(+)의 영향을 미칠 것이다.

H5: 과업-기술 적합성은 지속 사용 의도에 정(+)의 영향을 미칠 것이다.

3.2.3 사용자 만족, 과업 성과, 지속 사용 의도

사용자 만족은 생성형 AI를 활용한 사용자 경험에 기반한 긍정적인 평가를 의미한다. 정보시스템 분야 연구에서 만족은 사용자가 기술을 사용하기 이전에 가졌던 기대와 실제 사용 경험 간의 비교를 바탕으로 형성되는 주관적인 평가로 정의되며, 이는 정보 기술에 대한 전반적인 사용 경험과 성능을 반영하는 핵심 개념으로 간주된다(DeLone and McLean, 2003; Zhou and Ma, 2005). 이러한 사용자 만족은 정보기술 활용 수준을 결정하는 데 있어 중요

한 요인으로 작용해 왔다. 즉, 사용자의 만족 수준이 높을수록 해당 기술을 지속적으로 사용할 가능성이 높아지고, 다양한 과업 수행 상황에서 적극적으로 활용할 가능성도 증가하는 것으로 나타났다(Jo, 2024). 나아가, 기술에 대해 높은 만족감을 느낄수록 개인의 과업 성과에도 직접적인 긍정적인 영향을 미치는 것으로 보고되고 있다(정승민, 2024). 이에 본 연구는 생성형 AI 사용 경험에서 발생한 사용자 만족이 과업 성과 및 지속 사용 의도에 어떠한 영향을 미치는지를 실증적으로 검증하고자 하며, 이를 위해 다음과 같은 연구가설(H6~H7)을 설정하였다.

H6: 사용자 만족은 과업성과에 정(+)의 영향을 미칠 것이다.

H7: 사용자 만족은 지속 사용 의도에 정(+)의 영향을 미칠 것이다.

IV. 연구방법 및 실증분석

4.1 측정도구 개발 및 표본

본 연구는 SOR 이론적 프레임워크를 기반으로, 생성형 AI를 활용한 과업성과 및 지속 사용 의도에 대한 영향을 분석하기 위해 과업 수행 목적으로 ChatGPT를 사용하는 국내 사용자를 연구 대상으로 설정하였다. 제시된 연구모형의 각 변수를 측정하기 위해 응답자들은 모든 문항에 대해 (1) 강한 부정(전혀 그렇지 않다)에서 (5) 강한 긍정(매우 그렇다)에 이르는 5점 리커트 척도(5-point Likert scale)를 사용하여 응답하였다. 각 변수의 측정을 위한 설문 문항은 총 세 단계를 거쳐 개발되었다. 첫째, 선행연구를

바탕으로 각 잠재변수를 측정할 수 있는 초기 문항들을 수집한 후 본 연구의 맥락에 맞게 수정 및 보완하였다.

둘째, 관련 분야 연구자(예, 교수, 연구원, 대학원생 등)를 대상으로 내용 타당성(content validity)을 검토하여 문항의 표현과 적합성을 정교화하였다. 셋째, 사전 조사(pre-test)를 통해 측정모형의 신뢰성과 타당성을 통계적으로 검증하였다($n=85$). 사전 조사 결과, 신뢰성과 타당성을 저해하는 문항은 발견되지 않았으며, 이를 바탕으로 본 연구에서 사용할 최종 측정항목을 확정하였다. 다음의 <표 1>은 본 연구모형의 구성요소와 각 변수에 대한 조작적 정의 그리고 관련 선행연구를 요약하여 제시한 것이다.

자료 수집은 2025년 4월부터 5월까지 약 2개월간, 학습 또는 업무 목적으로 ChatGPT를 활용하는 국내 대학생과 일반 직장인을 대상으로 실시되었다. 총 327부의 설문 응답이 수집되었으며, 이 중 응답이 불완전하거나 무성의한 12부를 제외한 315부의 응답이 최종 분석에 사용되었다.

응답자의 인구통계학적 특성은 <표 2>에 요약되어 있으며, 주요 내용은 다음과 같다. 먼저 성별 분포는 남성 148명(46.98%), 여성 167명(53.02%)으로 성별 간 큰 편차 없이 비교적 균형 있게 나타났다. 이는 ChatGPT 사용이 성별에 관계없이 광범위하게 확산되고 있음을 시사한다. 연령대는 20대와 30대가 전체 응답자의 70%(285명)를 차지하여 젊은 연령층에서의 사용 비율이 특히 높은 것으로 나타났다. 이는 40대 이상에 비해 상대적으로 젊은 세대가 생성형 AI를 더 적극적으로 활용하고 있음을 보여준다.

ChatGPT 사용 기간은 3개월~6개월 미만인 125명(39.68%)으로 가장 많았으며 이어 6개월~1년 미만 96명(30.48%), 1개월~3개월 미만 46명(14.60%), 1년 이상 30명(9.52%), 1개월 미만 18명(5.71%)

〈표 1〉 연구변수에 대한 조작적 정의 및 관련연구

연구 변수	조작적 정의	관련연구
	측정항목	
의인화	사용자가 ChatGPT에 대해 인간과 유사한 행동양식, 감정, 동기 등이 부여되어 있다고 인식하는 정도	Gupta and Mukherjee(2024) Zhou and Ma (2025)
	① ChatGPT와의 대화는 실제 인간과 소통하는 것처럼 느껴진다. ② ChatGPT는 마치 인간과 같은 인식 능력을 가진 것처럼 느껴진다. ③ ChatGPT는 대화하는 상대방을 이해하는 듯한 느낌을 제공해 준다.	
개인화	ChatGPT가 제공하는 정보와 콘텐츠가 사용자 개인의 선호와 필요에 맞춰져 있다고 인식하는 정도	Zhou et al. (2022) Zhou and Ma (2025)
	① ChatGPT는 나에게 맞춤형 정보와 지식을 제공해 준다. ② ChatGPT에서 얻는 정보와 지식은 내 개인적인 관심사에 맞춰져 있다. ③ ChatGPT는 내 개인정보를 활용해 적절한 정보와 지식을 제공해 준다.	
지능성	ChatGPT를 사용할 때 똑똑하고 유능하다고 평가하는 주관적인 인식의 정도	Gupta and Mukherjee(2024) Zhou and Ma (2025)
	① ChatGPT는 내가 원하는 작업을 빠르게 처리하는 데 도움을 준다. ② ChatGPT는 내가 이해할 수 있는 방식으로 소통할 수 있다. ③ ChatGPT는 내가 내리는 지시를 이해할 수 있다.	
최소성	ChatGPT가 제공하는 정보나 콘텐츠가 다른 AI시스템이나 일반적인 정보 플랫폼에서 쉽게 얻을 수 없는 독창성과 차별성을 지니고 있다고 인식하는 정도	Zhou et al. (2022)
	① ChatGPT에서 생성한 정보는 다른 곳에서는 쉽게 얻기 어렵다. ② ChatGPT에서 제공하는 정보는 다른 플랫폼과의 정보와는 다르다. ③ ChatGPT에서 얻는 정보는 다른 사람에게서 쉽게 얻을 수 없다.	
정확성	ChatGPT가 제공하는 정보와 콘텐츠가 사용자의 요구와 일치하며 정확한 정보로 인식되는 정도	Gupta(2025) Zhou and Ma (2025)
	① ChatGPT가 제공하는 정보는 정확하다. ② ChatGPT가 제공하는 정보는 최신이다. ③ ChatGPT가 제공하는 정보는 믿을 수 있다.	
사용자 만족	ChatGPT를 사용하기 이전에 가졌던 기대와 실제 사용 경험 간의 비교를 바탕으로 형성되는 긍정적인 평가의 정도	Jo(2024) Zhou and Ma (2025)
	① ChatGPT는 내가 기대한 것보다 더 나은 경험을 제공한다. ② ChatGPT는 내 기대를 충족시킨다. ③ ChatGPT는 내가 원하는 정보나 지식을 잘 제공해 준다. ④ 전반적으로 나는 ChatGPT 사용 경험에 만족한다.	
과업-기술 적합성	ChatGPT가 업무를 수행하는데 어느 정도 일치하는가에 대해 인식하는 정도	Chakraborty et al. (2025) Zhang et al. (2025)
	① ChatGPT는 업무(학습) 수행에 필요한 정보와 지식을 제공해 준다. ② 내 업무(학습)에 ChatGPT를 사용하는 것은 매우 유용하다. ③ ChatGPT의 기능은 내가 원하는 목적에 매우 잘 맞는다. ④ ChatGPT는 내가 업무(학습)의 요구사항에 효과적으로 대응할 수 있다.	
과업성과	ChatGPT가 업무 수행을 지원하는 정도	DeLone and McLean (2003)
	① ChatGPT는 내가 업무(학습)의 정확성을 향상해 준다. ② ChatGPT를 사용하면 업무(학습)를 더 효율적으로 수행할 수 있다. ③ ChatGPT를 사용하는 것은 내 업무(학습)의 생산성을 높여준다. ④ ChatGPT를 사용하는 것은 내 업무(학습)의 성과를 향상하는 데 도움이 된다.	
지속 사용 의도	ChatGPT를 지속적으로 사용할 의지의 정도	Duong et al. (2024) Lam(2025) Jo(2024)
	① 나는 앞으로도 ChatGPT를 계속 사용할 것이다. ② 나는 자주 ChatGPT를 사용할 것이다. ③ 나는 다른 사람들에게 ChatGPT를 추천할 의향이 있다. ④ 나는 업무(학습)에 필요할 때 항상 ChatGPT를 사용할 계획이다.	

〈표 2〉 응답자의 특성

분류		빈도	응답비율
성별	남자	148	46.98%
	여자	167	53.02%
연령	20-29세	118	37.46%
	30-39세	109	34.60%
	40-49세	58	18.41%
	50세 이상	30	9.52%
최종 학력	고졸	51	16.19%
	대졸(전문대 포함)	138	43.81%
	대학원졸	126	40.00%
ChatGPT 사용기간	1개월 미만	18	5.71%
	1개월 ~ 3개월 미만	46	14.60%
	3개월 ~ 6개월 미만	125	39.68%
	6개월 ~ 1년 미만	96	30.48%
	1년 이상	30	9.52%
ChatGPT 사용 방식	무료 서비스 이용	189	60.00%
	유료 서비스 이용	126	40.00%
ChatGPT 사용 목적 (복수응답)	정보탐색 및 지식 습득(예: 개념 설명, 뉴스, 학술정보 등)	143	45.40%
	학습 또는 업무 과제 도움(예: 요약, 번역, 리포트 작성 등)	177	56.19%
	업무 지원(예: 이메일 작성, 보고서 초안, 기획 아이디어 등)	139	44.13%
	코딩 및 기술 개발(예: 프로그래밍 코드 생성, 오류 수정 등)	85	26.98%
	글쓰기 및 창작 활동(예: 에세이, 블로그, 소설 등)	181	57.46%
	일상 대화 및 심리적 위로	60	19.05%
	이미지 생성 및 디자인 아이디어 구상	93	29.52%
	기타	17	5.40%
합계		315	100%

순으로 나타났다. 이는 많은 사용자가 최근 1년 이내에 ChatGPT를 경험한 초기 사용자임을 시사한다. 서비스 이용 형태로는 무료 서비스 이용자가 189명(60.0%)으로 다수를 차지했으며, 유료 서비스 이용자는 126명(40.0%)으로 나타났다. 이는 여전히 무료 접근성을 기반으로 사용이 확대되고 있음을 보여준다. 마지막으로 ChatGPT의 사용 목적

(복수 응답)으로는 ‘글쓰기 및 창작 활동’이 181명(57.46%)으로 가장 많았으며, 다음으로 ‘학습 또는 업무 과제 해결’(177명, 56.19%), ‘정보 탐색 및 지식습득’(143명, 45.40%), ‘업무 지원’(139명, 44.13%) 등의 순으로 나타났다. 이는 ChatGPT가 창의적 활동과 과제 수행, 정보 활용 등 다양한 목적으로 폭넓게 사용되고 있음을 보여준다.

4.2 측정모형 검증결과

본 연구에서 측정도구의 적절성을 평가하기 위해 모형 적합도 검증, 내적일관성(신뢰성) 검증 그리고 타당성 검증(집중타당성과 판별타당성)을 실시하였다. 이를 위해 확인적 요인분석(Confirmatory Factor Analysis, CFA)을 수행하였으며, 분석에는 AMOS 29.0을 사용하였다. CFA 분석 결과를 바탕으로 측정모형의 적합도 지수 및 요인적재값을 통해 적합도와 집중타당성을 평가하였다. 모형 적합도 평가는 기존 연구에서 널리 활용되는 적합도 지수인 규범적합지수(NFI), 비교부합지수(CFI), 기초부합지수(GFI), 수정된 기초부합지수(AGFI), 상대적 카이제곱값(χ^2/df), 표준적합지수(RMSEA)를 사용하였다.

초기 CFA 분석 결과, 대부분의 지수가 기준치를 충족하였으나 NFI와 RMSEA는 각각 임계치인 0.9 이상과 0.05 이하에 충족하지 않는 것으로 나타났다. 이에 따라 CFA를 통해 산출되는 수정지수(Modification Indices, M.I)를 검토한 결과, 사용자 만족을 측정하는 첫 번째 항목(sat1)이 해당 잠재변수 외에 다른 잠재변수와도 연관성이 높은 것으로 나타났다(M.I = 21.83). 보수적인 기준에서 MI 값을 10으로 봤을 때, 산출값이 그 이상이라는 것은 이 문항이 단일차원성 위반 가능성의 문제가 있음을 의미한다. 즉, sat1 문항이 사용자 만족 외에 다른 요인과도 높게 상관을 가지는 교차 상관관계(cross-correlation)가 있다는 것을 시사한다. 따라서 sat1

문항을 제거한 후 CFA 검증을 다시 실시하였으며, 그 결과 모든 적합도 지수들이 임계치를 충족하는 것으로 나타나 측정모형의 전반적인 적합도는 확보된 것으로 판단되었다. 최종 모형의 적합도 지수는 <표 3>에 제시되어 있다.

구조모형 검토에 앞서, 최종 수집된 응답 자료(n = 315)를 바탕으로 측정도구의 신뢰성과 타당성을 검토하였다. 신뢰성 검증은 측정 항목 간의 내적 일관성(internal consistency)을 평가하기 위해 실증연구에서 가장 널리 사용되는 Cronbach's Alpha 계수를 사용하였다(Nunnally, 1967). 일반적으로 Cronbach's Alpha 값이 0.7 이상이면 신뢰성이 확보된 것으로 간주된다. 분석 결과, Cronbach's Alpha 값은 0.751에서 0.911로 나타나 측정도구의 신뢰성은 확보된 것으로 판단된다.

다음으로 타당성 검증은 집중타당성(convergent validity)과 판별타당성(discriminant validity)으로 구분하여 실시하였다. 집중타당성 검증을 위해 AMOS 29.0을 이용한 확인적 요인분석 결과로부터 요인적재값(factor loading), 합성신뢰도(Composite Reliability, CR) 및 평균분산추출(Average Variance Extracted, AVE) 값을 사용하였다. 일반적으로 요인적재량은 ± 0.4 이상이면 유의한 것으로 판단되며(Barclay et al., 1995), 합성신뢰도는 0.7 이상, 각 잠재변수의 AVE 값은 0.5 이상일 때 집중타당성이 확보된 것으로 판단한다(Fornell and Lacker, 1981). 또한, 탐색적요인분석(Exploratory Factor

<표 3> 측정모형의 적합도 검증

모델	NFI	GFI	AGFI	CFI	χ^2/df	RMSEA
초기측정모형	0.876	0.911	0.829	0.910	2.883	0.079
수정된 측정모형	0.918	0.927	0.880	0.935	2.005	0.049
권장치	≥ 0.9	≥ 0.9	≥ 0.8	≥ 0.9	≤ 3.0	≤ 0.05

〈표 4〉 측정변수의 신뢰성 및 타당성 분석 결과

변수	항목	평균	S.D	요인값	C.R	Cronbach's Alpha	합성신뢰도	AVE
의인화 (Anthropomorphism)	ant1	5.141	0.758	0.764	-	0.843	0.844	0.643
	ant2	5.075	0.738	0.801	19.657			
	ant3	5.339	0.852	0.839	19.248			
개인화 (Personalization)	per1	6.089	0.578	0.833	-	0.832	0.849	0.670
	per2	6.315	0.609	0.827	15.783			
	per3	5.329	0.645	0.796	19.100			
지능성 (Intelligence)	int1	5.046	0.768	0.805	-	0.751	0.850	0.654
	int2	5.527	0.631	0.798	18.425			
	int3	5.758	0.631	0.822	17.867			
희소성 (Rareness)	rar1	6.466	0.610	0.800	-	0.789	0.862	0.675
	rar2	4.683	0.744	0.823	18.100			
	rar3	5.716	0.550	0.842	18.692			
정확성 (Accuracy)	acc1	4.542	0.850	0.877	-	0.804	0.896	0.741
	acc2	5.542	0.622	0.863	20.047			
	acc3	5.626	0.698	0.842	18.232			
사용자 만족 (User Satisfaction)	sat1	삭제				0.886	0.915	0.783
	sat2	5.957	0.693	0.914	-			
	sat3	5.872	0.531	0.887	16.880			
	sat4	5.928	0.565	0.853	17.537			
과업-기술 적합성 (Task-Technology Fit)	tf1	5.806	0.566	0.841	-	0.900	0.925	0.756
	tf2	5.740	0.567	0.809	19.540			
	tf3	5.730	0.516	0.897	17.774			
	tf4	5.867	0.604	0.926	20.258			
과업 성과 (Task Performance)	tp1	5.504	1.338	0.833	-	0.827	0.920	0.742
	tp2	5.131	1.205	0.875	18.711			
	tp3	4.296	1.080	0.910	17.824			
	tp4	3.867	0.970	0.825	19.456			
지속 사용 의도 (Intention to Continue Use)	icu1	6.230	0.706	0.899	-	0.911	0.945	0.810
	icu2	6.301	0.533	0.897	18.314			
	icu3	5.707	0.749	0.890	17.543			
	icu4	5.884	0.367	0.914	16.852			

주) “-”: 분석시 “1”로 고정함.

〈표 5〉 잠재변수의 판별타당성 분석결과

변수	1	2	3	4	5	6	7	8	9
1. 의인화	0.802								
2. 개인화	0.324	0.819							
3. 지능성	0.219	0.331	0.808						
4. 회소성	0.309	0.474	0.232	0.822					
5. 정확성	0.263	0.379	0.291	0.324	0.861				
6. 사용자 만족	0.225	0.187	0.200	0.382	0.401	0.885			
7. 과업-기술 적합성	0.233	0.316	0.265	0.323	0.216	0.208	0.869		
8. 과업 성과	0.206	0.209	0.338	0.450	0.311	0.261	0.310	0.861	
9. 지속 사용 의도	0.214	0.323	0.259	0.198	0.284	0.362	0.368	0.248	0.900

주) 진하게 표시된 대각선 AVE의 제곱근 값임.

Analysis: EFA)을 실시해 측정항목의 타당성에 문제가 없음을 다시 확인하였다. EFA 분석 결과는 부록 1에서 보여주고 있으며, 분석 결과 각 측정항목의 요인적재값은 권장치 이상으로 나타났으며, 교차요인에 대한 이슈는 없는 것으로 확인되었다.

판별타당성 검정은 Fornell and Larcker(1981)가 제시한 평균분산추출과 Pearson 상관관계분석 방법을 사용하였다. 판별타당성 존재 여부에 대한 검증은 각 잠재변수의 AVE의 제곱근 값이 해당 변수와 다른 변수 간의 Pearson 상관계수보다 클 경우 판별타당성이 확보된 것으로 간주한다. 검증 결과, 모든 잠재변수의 AVE 제곱근이 종과 횡의 상관계수보다 높은 값을 보여, 측정모형의 판별타당성 역시 문제가 없는 것으로 확인되었다. 이와 같은 분석 결과는 본 연구의 설문 문항들이 통계적으로 유의한 수준의 신뢰성과 타당성을 갖추고 있음을 증명한다. 분석 결과는 〈표 4〉와 〈표 5〉에 제시되어 있다.

마지막으로 본 연구와 같이 설문조사에서 동일한 측정으로 야기 될 수 있는 공통방법편의(Common Method Bias: CMV)의 가능성을 검토하기 위해 Harman의 단일요인 검정(Harman's single-factor

test)을 실시하였다. 주성분 분석(Principal Component Analysis)을 활용한 비회전 요인분석 결과, 총 9개의 요인이 추출되었으며, 이 중 첫 번째 요인이 설명하는 분산은 31.4%로 이는 전체 분산의 과반수에 미치지 않는 수치이다. 이는 Malhotra et al.(2007)이 제시한 특정 요인에 대한 분산의 40% 기준과도 부합한다. 따라서 본 연구에서의 CMB는 주요한 문제 요인으로 작용하지 않을 가능성이 크다고 판단된다(부록 2 참조.).

4.3 구조모형 검증결과

측정모형에 대한 신뢰성과 타당성 검정 후, 본 연구에서는 제안된 연구모형 내 변수 간의 인과관계를 실증적으로 분석하기 위해 구조방정식모형(Structural Equation Modeling, SEM) 기법을 활용하였다. 구조방정식모형 분석을 통해 모형에 대한 전반적인 적합도뿐만 아니라 잠재변수 간의 경로 관계와 내생 변수의 설명력을 나타내는 결정계수(R^2)에 대해서도 확인할 수 있다.

구조모형의 적합도 검정 결과는 측정모형의 적합

도 검정에서 사용한 것과 동일한 지수들을 기준으로 평가하였다. 분석 결과, 규범적합지수(NFI)는 0.924, 기초부합지수(GFI)는 0.930, 수정된 기초부합지수(AGFI)는 0.897, 비교부합지수(CFI)는 0.939, 상대적 카이제곱 값(χ^2/df)은 2.001, 표준적합지수(RMSEA)는 0.037로 나타났으며, 이들 지수 모두 일반적으로 제시되는 기준치를 충족하는 것으로 나타나 구조모형의 적합도는 확보한 것으로 판단되었다.

생성형 AI의 다섯 가지 특성이 사용자 만족에 미치는 영향을 검증한 결과를 살펴보면, 의인화 경로계수(β)는 0.429이며, t값은 5.034로 유의수준 0.01에서 통계적으로 유의하였다. 개인화는 경로계수 0.350, t값 4.827로 유의수준 0.01에서 유의하였으며, 지능성은 경로계수 0.463, t값 9.520, 회소성은 경로계수 0.266, t값 3.891, 정확성은 경로계수 0.390, t값 6.632로 모두 유의수준 0.01에서 통계적으로 유

의하였다. 따라서 H1a부터 H1e까지 모든 가설이 채택되었다.

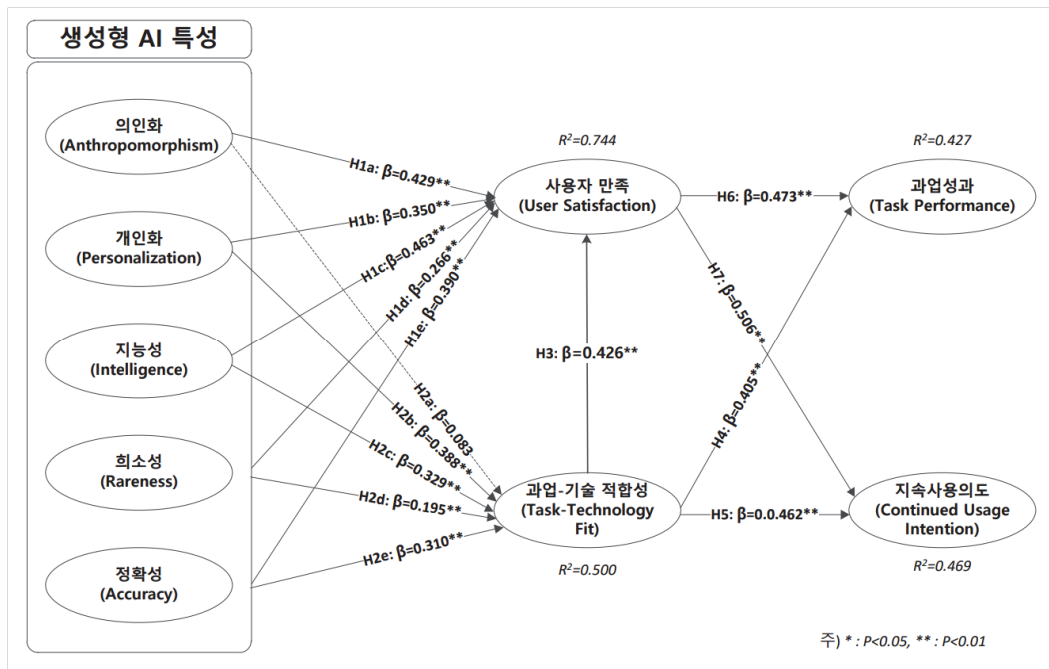
다음으로 생성형 AI 특성이 과업-기술 적합성에 미치는 영향을 분석한 결과, 개인화는 경로계수 0.388, t값 6.073으로 유의수준 0.01에서 유의하였으며, 지능성은 경로계수 0.329, t값 4.753으로 유의하였다. 정확성은 경로계수 0.310, t값 4.902로 유의수준 0.01에서 유의하게 나타났으며, 회소성은 경로계수 0.195, t값 2.323로 유의수준 0.05에서 유의성을 확보하였다. 반면 의인화는 통계적으로 유의하지 않는 것으로 나타나 H2a는 기각되었고, H2b부터 H2e까지는 모두 채택되었다.

과업-기술 적합성이 사용자 만족에 미치는 영향을 분석한 결과, 과업-기술 적합성은 경로계수 0.426, t값 5.951로 유의수준 0.01에서 유의하게 나타났으며 이에 따라 H3은 채택되었다. 또한, 과업-기술 적

〈표 6〉 가설검정 결과요약

가설	경로			경로계수	t-값	채택 유·무
H1a	의인화	→	사용자 만족	0.429**	5.034	채택
H1b	개인화			0.350**	4.827	채택
H1c	지능성			0.463**	9.520	채택
H1d	회소성			0.266**	3.891	채택
H1e	정확성			0.390**	6.632	채택
H2a	의인화	→	과업-기술 적합성	0.083	1.027	기각
H2b	개인화			0.388**	6.073	채택
H2c	지능성			0.329**	4.753	채택
H2d	회소성			0.195*	2.323	채택
H2e	정확성			0.310**	4.902	채택
H3	과업-기술 적합성	→	사용자 만족	0.426**	5.951	채택
H4			과업성과	0.405**	6.227	채택
H5			지속 사용 의도	0.462**	8.990	채택
H6	사용자 만족	→	과업성과	0.473**	8.319	채택
H7			지속 사용 의도	0.506**	10.258	채택

주) *: $p < 0.05$, **: $p < 0.01$, R^2 (사용자 만족) = 0.744, R^2 (과업-기술 적합성) = 0.500, R^2 (과업성과) = 0.427, R^2 (지속 사용 의도) = 0.469.



〈그림 2〉 구조모형 분석결과

합성은 과업 성과에 대해 경로계수 0.405, t 값 6.227로, 지속 사용 의도에 대해서는 경로계수 0.462, t 값 8.990으로 유의수준 0.01에서 유의한 영향을 미치는 것으로 나타났다. 따라서 H3부터 H5까지의 가설은 모두 채택되었다. 사용자 만족이 과업 성과와 지속 사용 의도에 미치는 영향을 분석한 결과에서는 사용자 만족은 과업 성과에 대해 경로계수 0.473, t 값 8.319로 나타났으며, 지속 사용 의도에 대해서는 경로계수 0.506, t 값 10.258을 보여 유의수준 0.01에서 유의한 영향을 미치는 것으로 나타났다. 이에 따라 H6과 H7 또한 채택되었다.

한편, 구조모형 분석을 통해 내생변수에 대한 결정계수(R^2)를 확인한 결과, 사용자 만족의 R^2 값은 0.744로 나타나 외생변수들이 전체 분산의 약 74.4%를 설명하고 있는 것으로 확인되었다. 과업-기술 적합

성의 R^2 값은 0.500, 과업성과의 R^2 값은 0.427, 지속 사용 의도의 R^2 값은 0.469로 나타나 본 연구 모형이 각 내생변수에 대해 중간 이상의 설명력을 갖추고 있는 것으로 평가된다. 구조모형에 대한 실증분석 결과는 〈표 6〉과 〈그림 2〉에 제시되어 있다.

V. 결 론

5.1 연구결과와 시사점

ChatGPT를 비롯한 생성형 AI가 출시된 지 3년이 채 되지 않았음에도 불구하고, 대화형 기반의 생성형 AI는 개인과 조직 전반에 혁신적인 변화를 불러오고

있다. 단순한 정보 검색이나 데이터 처리방식에 머물렀던 기존의 정보기술들과 다르게 생성형 AI는 자연어 처리 기반의 고도화된 상호작용을 통해 보다 직관적이고 창의적인 방식의 문제 해결을 가능하게 하고 있다. 조직에서도 생성형 AI는 단순히 업무의 효율성을 높이는 도구로 활용되는 것을 넘어서 AI 기술을 중심으로 업무 구조와 방식을 근본적으로 재설계하려는 움직임도 활발히 나타나고 있다. 개인 또한 생성형 AI를 단순한 정보 탐색 도구로 인식하는 것을 넘어, 사용자의 요구에 따라 자율적으로 행동을 수행할 수 있는 능동적 기술로 AI를 받아들이고 있다. 이러한 인식의 변화는 생성형 AI를 일상생활, 의사결정, 학습, 창작 등 다양한 개인 활동에 적극적으로 활용하는 흐름으로 이어지고 있으며 이는 정보 기술 수용에 대한 개인의 행동과 디지털 역량 전반에 새로운 전환점을 마련하고 있다.

이에 본 연구에서는 생성형 AI에 대한 개인의 사용 경험과 행동을 실증적으로 분석하고자 하였다. 이를 위해 환경적 요인이 인간의 의지와 행동에 미치는 영향을 설명하는 데 유용한 이론적 틀로 알려져 있는 S-O-R 프레임워크를 연구모형의 기반으로 활용하였다. 나아가, 생성형 AI가 개인 및 조직 차원에서 다양한 과업 수행에 활용되고 있다는 점을 고려하여 과업과 기술 간의 적합성이 사용자 행동과 성과에 미치는 영향을 설명하는 이론인 TTF를 보완적으로 사용하였다. 이러한 이론적 문헌 고찰을 기반으로 본 연구는 생성형 AI를 대표하는 서비스인 ChatGPT 사용자들을 대상으로 실증분석을 수행하였다. S-O-R 프레임워크의 자극-유기체-반응 구조에 따라 ChatGPT의 특성, 사용자 만족, 과업-기술 적합성, 과업 성과와 지속 사용 의도를 주요 변수로 설정하여 이들 간의 관계를 분석하기 위한 연구모형을 구성하여 가설검증을 진행하였다. 이에 수집된 자료를 분석한

결과 요약과 이에 대한 토의는 다음과 같다.

첫째, 생성형 AI의 특성으로 제안한 개인화, 지능성, 희소성, 정확성은 사용자 만족과 과업-기술 적합성에 유의한 영향을 미치는 것으로 나타났다. 이는 사용자가 생성형 AI와 상호작용하는 과정에서 해당 기술이 사용자의 요구에 맞춰 콘텐츠와 정보를 신속하게 제공하고, 사용자가 이해하기 쉬운 방식으로 의사소통할 수 있다고 인식할수록 만족도가 높아짐을 시사한다. 또한, 생성형 AI가 제공하는 콘텐츠나 정보가 다른 플랫폼에서는 쉽게 얻기 어렵고 차별화된 가치를 가지며 정확하고 최신의 정보를 기반으로 한다고 인식할수록 사용자 만족에 긍정적인 영향을 미침을 의미한다. 나아가, 사용자가 생성형 AI와의 상호작용을 통해 과업 수행에 실질적인 도움을 받을 수 있다고 판단하고, 생성형 AI 특성이 과업 수행에 유용하다고 인식할수록 해당 기술이 자신의 과업에 적합한 도구라고 인식하게 되며 과업-기술 적합성에 대한 인식 향상으로 이어짐을 보여주고 있다. 하지만, 생성형 AI의 특성 중 하나로 제안한 의인화는 사용자 만족에는 유의한 영향을 미치는 것으로 나타났으나 과업-기술 적합성에는 유의한 영향을 미치지 않는 것으로 나타났다. 이는 사용자가 생성형 AI와의 상호작용을 통해 마치 인간과 대화하는 것 같은 느낌을 받거나 기술이 자신의 의도를 이해하는 듯한 경험을 할 경우 만족도는 높아질 수 있음을 시사한다. 그러나 사용자는 이러한 인간 유사성에 대한 인식이 실제 과업 수행에 필요한 기술적 적합성으로는 연결되지 않는다고 인식하고 있음을 알 수 있다.

둘째, 과업-기술 적합성은 사용자 만족, 과업 성과, 지속 사용 의도에 모두 긍정적인 영향을 미치는 것으로 나타났다. 이러한 결과는 사용자가 생성형 AI를 과업 수행을 지원하는 적합한 기술로 인식할수록 전반적인 만족도가 높아지고, 과업을 보다 효율적으로

수행할 수 있음을 시사한다. 또한, 생성형 AI가 과업 수행에 유용한 기술이라고 인식할수록 사용자가 해당 기술을 지속적으로 사용하려는 의지로 이어질 수 있음을 보여준다.

셋째, 사용자 만족은 과업 성과와 지속 사용 의도에 긍정적인 영향을 미치는 것으로 나타났다. 이는 생성형 AI를 사용하는 과정에서 느끼는 만족도가 높을수록 자신의 과업에서 더 나은 성과를 낼 수 있으며, 해당 기술을 지속적으로 사용할 가능성이 높아진다는 것을 의미한다.

5.2 연구의 시사점

본 연구는 생성형 AI 활용과 관련하여 다음과 같은 이론적 시사점을 제공한다. 첫째, 본 연구는 생성형 AI가 학습, 업무, 문제 해결, 창작 등 다양한 실제 과업 수행에 활용되는 사례가 증가하고 있는 시점에서 생성형 AI 사용자의 지속적인 사용 행동과 과업 성과에 주목하였다는 점에서 의의가 있다. 특히, 생성형 AI는 단순한 정보 검색 플랫폼이나 기존 AI와는 본질적으로 다른 특성들을 지니고 있음에도 불구하고 이를 고려한 사용자 행동에 관한 연구가 아직은 초기 단계에 머물러 있는 상황에서 차별화되는 생성형 AI 특성을 선행연구의 이론적 고찰을 통해 도출하고 그 관계를 실증적으로 검증함으로써 생성형 AI 연구의 범위를 확장했다는 데 이론적 의의가 있다.

둘째, 이론적 틀의 측면에서 본 연구는 기존 연구들이 주로 단일 이론을 중심으로 생성형 AI 수용 행동을 설명하고 있는 것에서 한 단계 나아가, S-O-R 프레임워크와 TTF 이론을 통합적으로 고려했다는 점에서 의의가 있다. S-O-R 프레임워크는 인간 행동의 원인을 파악하고 관련 문제를 이해하는 데 유용한 이론으로 본 연구는 이를 기반으로 생성형 AI의 특성

(외부 자극)이 사용자 만족(유기체)과 과업 성과 및 지속 사용 의도(반응)로 이어지는 구조를 제안하였다. 또한, TTF 이론의 과업-기술 적합성을 유기체 요인으로 함께 제안하여 그 관계를 실증분석을 통해 검증함으로써 기존 연구의 설명력을 보완하는 통합적인 이론적 틀을 제시하였다. 특히, 개인 사용자의 학습, 업무 등의 환경에서 생성형 AI의 활용이 증가하고 있는 상황에서 기존에는 실무 시스템의 활용에 관한 연구에서 주로 적용되었던 TTF 이론을 개인 사용자 환경에 적용하였다는 점에서 이론적 의의가 있다.

셋째, 본 연구는 생성형 AI의 특성을 단순한 정보 기술의 기술적 특성이 아니라 사용자가 생성형 AI를 사용하는 과정에서 인지하는 개인화, 의인화, 지능성, 희소성, 정확성과 같은 기존 서비스와의 차별화되는 요인으로 구조화하고 변수화하였다는 점에서 의의가 있다. 이는 기존 연구들이 생성형 AI를 기술 그 자체로 다루는데 그쳤던 반면, 본 연구는 사용자 지각을 중심으로 이론적 틀을 제시하고 개념적 정의와 조작화를 시도하였다는 데에서 이론적 의미가 있다.

본 연구의 실무적 시사점은 다음과 같다. 첫째, 본 연구는 생성형 AI 사용자 경험을 구조모형으로 분석하여 사용자 만족과 과업-기술 적합성이 과업성과 지속사용의도에 유의미한 영향을 미친다는 점을 실증적으로 밝혔다. 이러한 결과는 생성형 AI 개발자와 플랫폼 운영자에게 기술 개발 및 운영 전략의 방향성을 제시한다. 특히, 생성형 AI가 기존의 정보 검색 플랫폼이나 전통적 AI 시스템과 차별화되는 특성으로 개인화, 지능성, 희소성, 정확성 등의 요소가 중요한 영향을 미친다는 점은 사용자 중심의 경험을 강화하기 위한 기술 설계의 필요성을 시사한다. 또한, 생성형 AI의 한계점으로 지적되고 있는 AI 환각(hallucination) 문제를 최소화하고 실제 과업 수행에 실질적인 도움을 줄 수 있도록 과업 지향적 응답

품질을 지속해서 개선할 필요가 있다. 이를 통해 사용자 만족을 확보하고 장기적인 플랫폼 이용을 유도할 수 있을 것이다.

둘째, 본 연구의 결과는 일반 사용자 및 실무자에게도 중요한 시사점을 제공한다. 사용자 만족이 과업 성과와 지속사용의도에 영향을 미친다는 점은 초기 사용 경험이 후속 이용 행태에 결정적인 영향을 미친다는 것을 의미한다. 따라서 생성형 AI를 활용하는 개인 사용자나 조직 내 실무자는 단순한 정보 조회를 넘어 다양한 과업에서 생성형 AI를 적극적으로 활용할 수 있는 방향으로 역량을 확대해 나갈 필요가 있다. 이를 위해서는 생성형 AI의 기능과 한계를 이해하고 적절한 활용 범위 내에서 과업을 지원받는 과정이 중요하다. 조직 차원에서도 초기 사용자가 생성형 AI에 대해 긍정적인 경험을 가질 수 있도록 적절한 사용 매뉴얼과 사전 교육 등을 통해 AI 활용 역량을 체계적으로 강화할 필요가 있다.

셋째, 생성형 AI의 도입 및 활용과 관련하여 기업 및 조직 차원의 전략적 활용 방안도 제안할 수 있다. 생성형 AI는 단순한 보조 도구를 넘어 실제 과업-기술 적합성이 높은 영역에서 전략적 기술 자원으로 활용될 수도 있다. 예를 들어, 스타벅스가 내부 운영 효율화를 위해 도입한 ‘그린닷 어시스트(Green Dot Assist)’ 플랫폼은 생성형 AI의 업무 지원 가능성을 보여주는 대표적 사례로 생성형 AI가 단순 반복 업무를 자동화하고 서비스 응답시간을 단축하며 전반적인 조직 운영의 효율성을 제고할 수 있음을 시사한다. 이에 따라 기업은 과업의 특성과 생성형 AI의 특성을 면밀하게 분석하고 적합한 업무 영역을 선정하여 점진적이면서도 전략적으로 기술을 도입하는 방식이 필요하다. 이는 업무 효율성과 조직 구성원 만족도 그리고 조직 전반의 생산성 제고에 긍정적인 과급효과를 가져올 수 있을 것이다.

넷째, 본 연구는 정책적 관점에서도 실무적 시사점을 제공한다. 생성형 AI의 과업성과에 대한 긍정적 영향은 교육이나 공공부문 등 다양한 분야에서 생성형 AI의 활용 가능성을 뒷받침한다. 특히, 고도의 정보처리나 민원 응대가 필요한 영역에서는 생성형 AI 도입을 통해 만족도 및 업무 효율성을 제고할 수 있다. 이에 따라 정부 및 관련 기관은 생성형 AI에 대한 사회 전반의 인식을 제고하고 실제 현장 적용을 위한 역량 강화 교육 및 활용 매뉴얼 제공 등 정책적 지원을 확대해야 한다. 또한, 생성형 AI의 활용 확대에 따라 예상되는 개인정보 침해, 허위 정보 생성, 알고리즘 편향 등과 같은 문제점에 대비한 제도적 장치 마련도 병행되어야 한다. 지속 가능한 AI 생태계 조성을 위해서는 기술의 발전과 함께 법적 제도, 윤리적 기준, 사회적 합의가 균형 있게 추진될 필요가 있다.

5.3 연구의 한계점 및 향후 연구 방향

본 연구는 ChatGPT 사용자를 대상으로 생성형 AI의 특성이 사용자 만족과 과업-기술 적합성을 매개로 과업 성과와 지속 사용 행동에 어떠한 영향을 미치는지를 실증적으로 분석함으로써 이론적 및 실무적으로 의미 있는 시사점을 도출하였다. 그럼에도 불구하고 본 연구는 다음과 같은 몇 가지 한계점을 지닌다. 먼저, 본 연구는 조사 시점을 기준으로 생성형 AI 서비스 중 가장 대중적으로 사용되고 있는 ChatGPT 사용자에 한정하여 실증분석을 수행하였으나 이는 생성형 AI 전반에 본 연구의 결과를 일반화하는 데 한계가 있을 수 있다. 따라서 향후 연구에서는 생성형 AI 플랫폼 시장점유율이나 사용자 선호도를 고려하여 ChatGPT 외의 다른 생성형 AI 또는 다양한 플랫폼 사용자로까지 연구 대상을 확대하여 본 연구모형을 적용해 볼 필요가 있다.

둘째, 본 연구는 생성형 AI와 사용자 간의 상호작용에 초점을 맞추었기 때문에 기존 AI 기술과 차별화되는 생성형 AI의 특성에 초점을 맞추어 연구를 진행하였다. 하지만 사용자의 신념이나 행동은 기술적 요인뿐만 아니라 다른 내·외부적 요인에 의해서도 영향을 받을 수 있다. 따라서 향후 연구에서는 생성형 AI 사용에 영향을 주는 여러 요인을 함께 고려하여 연구모형을 확장하거나 다양한 독립변수를 포함한 연구를 시도해 볼 필요가 있다. 또한, 과업-기술 적합성 측면에서 조직이나 개인의 환경적 요인(예: 조직 문화, 지원, 정책 등)도 함께 고려함으로써 생성형 AI의 실제 사용 맥락에 대한 이해를 더욱 심화시킬 수 있을 것이다. 나아가, 조절변수나 집단 간 차이에 초점을 두고 다양한 상황적 조건을 포함한 후속 연구를 통해 사용자의 인식이나 반응의 차이를 살펴보는 것도 유의미한 방향이 될 수 있을 것이다.

셋째, 본 연구는 과업 성과를 정성적 지표로 측정하고 있다. 이러한 방식은 사용자의 주관적 판단에 의존하기 때문에 실제 성과와 지각한 성과 간에 차이가 있을 가능성이 존재한다. 또한, 사용자의 과업 유형(예를 들어 개인적 과업 수행과 조직에서의 과업 수행)에 따라 과업 성과에 대한 응답에 차이가 있을 수도 있다. 따라서 향후 연구에서는 과업 성과를 보다 객관적인 성과지표나 정량적 지표를 활용하여 측정함으로써 분석의 신뢰도를 높일 필요가 있다. 마지막으로, 본 연구는 설문조사 방식에 기반하여 진행되었으나 향후 연구에서는 이러한 방법 외에도 생성형 AI 사용자와의 심층 인터뷰를 통해 질적 자료를 수집함으로써 사용자 행동의 이면을 보다 심층적으로 탐색할 필요가 있다.

참고문헌

- 김소연, “2년만에 연간 매출 13조5천억… 오픈AI 초고속 성장이끈 챗GPT,” <https://www.hankyung.com/article/2025061039707>, 2025년 6월 접속.
- (Kim, S. Y., “₩13.5 Trillion in Annual Revenue in Just Two Years: ChatGPT Drives OpenAI’s Explosive Growth,” <https://www.hankyung.com/article/2025061039707>, retrieved June 2025.)
- 김숙경 (2024), “생성형 AI의 성공적 도입 및 활용을 위한 전략적 접근,” **SW중심사회**, 제120권.
- (Kim, S. K. (2024). “A Strategic Approach for the Successful Adoption and Utilization of Generative AI,” *Journal of Software-Centric Society*, 120.)
- 박기호, 이군호 (2024). “ChatGPT 사용 만족도에 미치는 영향 요인 : 신뢰성의 매개효과,” **한국IT서비스학회지**, 제23권 2호, pp.99-116.
- (Park, K. H., and Li, J. H. (2024). “Factors Influencing User’s Satisfaction in ChatGPT Use: Mediating Effect of Reliability,” *Journal of Information Technology Services*, 23(2). pp.99-116.)
- 선담은, “오픈AI, 올해 매출 127억달러 예상… 지난해보다 3배 증가,” https://www.hani.co.kr/arti/economy/it/1189290.html?utm_source=copy&utm_medium=copy&utm_campaign=btn_share&utm_content=20250624, 2025년 6월 접속.
- (Sun, D. E. “OpenAI Projects \$12.7 Billion in Revenue This Year - Triple the Earnings from Last Year,” https://www.hani.co.kr/arti/economy/it/1189290.html?utm_source=copy&utm_medium=copy&utm_campaign=btn_share&utm_content=20250624, retrieved June 2025.)

- 선덕길, 유재현 (2024). “생성형 AI 서비스 사용자의 지속 사용의도에 미치는 영향: 가치기반수용모델을 중심으로,” **정보시스템연구**, 제33권 4호, pp.1-26.
- (Sun, D. G. and You, J. H. (2024). “A Study on the Factor Influencing User’s Intention to Continue Using Generative AI Services: Focusing on the Value-Based Adoption Model,” *The Journal of Information Systems*, 33(4), pp. 1-26.)
- 안승규, 안현철 (2024). “대화형 생성AI 서비스 사용자의 지속사용의도에 관한 연구: 과업-기술적합(TTF) 과 신뢰를 중심으로,” **경영정보학연구**, 제26권 1호, pp.193-218.
- (Ann, S. and Ahn, H. (2024). “A Study on User Continuance Intention of Conversational Generative AI Services: Focused on Task-Technology Fit (TTF) and Trust,” *Information Systems Review*, 26(1), pp.193-218.)
- 양지훈, 윤상혁 (2024). ChatGPT를 넘어 생성형(Generative) AI 시대로: 미디어 · 콘텐츠 생성형 AI 서비스사례와 경쟁력 확보 방안, **미디어이슈&트렌드**, 제55권, pp.62-70.
- (Yang, J. H. and Yoon, S. H. (2024). “Beyond ChatGPT: Generative AI in Media and Content - Service Cases and Competitive Strategies,” *Media Issue & Trend*, 55, pp. 62-70.)
- 이원준 (2024). “ChatGPT 같은 멘토, 우리 같은 멘티: AI 멘토쉽의 특성과 영향,” **경영학연구**, 제53권 6호, pp.1353-1374
- (Lee, W. J. (2024). “Mentor Like ChatGPT, Mentee Like Us: AI Mentorship’s Characteristics and Influence,” *Korean Management Review*, 53(6), pp.1353-1374.)
- 이진, 오현정 (2024). “생성형 AI의 기술적 특성과 사용자의 AI 리터러시가 생성형 AI의 지속사용의도에 미치는 영향,” **한국광고홍보학보**, 제26권 3호, pp. 96-140.
- (Lee, J. and Oh, H. J. (2024). “Determinants of Continuous Use of Generative AI: Focusing on the Extended Technology Acceptance Model (ETAM) and AI Literacy,” *The Korean Journal of Advertising and Public Relations*, 26(3), pp.96-140.)
- 이한신, 김판수 (2019). “소비자의 기술수용과 저항이 인공지능(AI) 사용의도에 미치는 영향,” **경영학연구**, 제48권 5호, pp. 1195-1219.
- (Lee, H. S. and Kim, P. S. (2019). “The Effect of Consumer’s Technology Acceptance and Resistance on Intention to Use of Artificial Intelligence (AI),” *Korean Management Review*, 48(5), pp. 1195-1219.)
- Abumalloh, R. A., Halabi, O., and Nilashi, M. (2025). “The Relationship Between Technology Trust and Behavioral Intention to Use Metaverse in Baby Monitoring Systems’ Design: Stimulus-Organism-Response (SOR) Theory,” *Entertainment Computing*, 52, pp.1-14.
- Albayati, H. (2024). “Investigating Undergraduate Students’ Perceptions and Awareness of Using ChatGPT as a Regular Assistance Tool: A User Acceptance Perspective Study,” *Computers and Education: Artificial Intelligence*, 6, pp.1-15.
- Barclay, D., R. Thompson, and C. Higgins (1995). “The Partial Least Squares(PLS) Approach to Casual Modeling: Personal Computer Adoption and Use as an Illustration,” *Technology Studies*, 2(2), pp.285-309.
- Belk, R. W. (1975). “Situational Variables and Consumer Behavior,” *Journal of Consumer Research*, 2(3), pp.157-164.
- Biswas, M. I., Talukder, M. S., and Chen, Y. (2025). “Applying the Stimulus-Organism Behavior-

- Consequence Framework to Examine the Relationship between Intention, Usage and Recommendation of ChatGPT in Higher Education," *International Journal of Educational Management*, 39(2), pp.450-468.
- Chakraborty, D., Troise, C., and Bresciani, S. (2025). "Exploring Consumer Intentions to Continue: Integrating Task Technology Fit and Social Technology Fit in Generative AI based Shopping Platforms," *Technovation*, 142, pp. 1-13.
- Cheng, X., Zhang, X., Cohen, J., and Mou, J. (2022). "Human vs. AI: Understanding the Impact of Anthropomorphism on Consumer Response to Chatbots from the Perspective of Trust and Relationship Norms," *Information Processing & Management*, 59(3), pp.1-16.
- Choi, J., Lee, H. J., and Kim, H. W. (2017). "Examining the Effects of Personalized App Recommender Systems on Purchase Intention: A Self and Social-Interaction PERSpective," *Journal of Electronic Commerce Research*, 18(1), pp. 73-102.
- Chung, M., Ko, E., Joung, H., and Kim, S.J. (2020). "Chatbot e-Service and Customer Satisfaction Regarding Luxury Brands," *Journal of Business Research*. 117, pp.587-595.
- DeLone, W. H. and McLean, E. R. (2003). "The DeLone and McLean Model of Information Systems Success: A Ten-Year Update," *Journal of Management Information Systems*, 19(4), pp.9-30.
- Duivenvoorde, B. (2025). "Generative AI and the Future of Marketing: A Consumer Protection Perspective," *Computer Law & Security Review: The International Journal of Technology Law and Practice*, 57, pp.1-14.
- Duong, C. D., Nguyen, T. H., Ngo, T. V. N., Dao, V. T., Do, N. D., and Pham, T. V. (2024). "Exploring Higher Education Students' Continuance Usage Intention of ChatGPT: Amalgamation of the Information System Success Model and the Stimulus-Organism-Response Paradigm," *The International Journal of Information and Learning Technology*, 41(5), pp.556-584.
- Fornell, C. and D. Larcker (1981). "Evaluating Structural Equation Models with Unobservable Variables and Measurement Error," *Journal of Marketing Research*, 18(1), pp.39-50.
- Goodhue, D. L. and Thompson, R. L. (1995). "Task-Technology Fit and Individual Performance," *MIS Quarterly*, 19(2), pp.213-236.
- Gupta, V. (2025). "Factors Influencing Librarian Adoption of ChatGPT Technology for Entrepreneurial Support," *Library Hi Tech*, ahead-of-print, pp.1-39.
- Gupta, R., Nar, K., Mishra, M., Ibrahim, B., and Bhardwaj, S. (2024). "Adoption and Impacts of Generative Artificial Intelligence: Theoretical Underpinnings and Research Agenda," *International Journal of Information Management Data Insights*, 4(1), pp.1-15.
- Gupta, A. S., and Mukherjee, J. (2024). "Framework for Adoption of Generative AI for Information Search of Retail Products and Services," *International Journal of Retail & Distribution Management*, 53(2), pp.165-181.
- Howard, M. C. and Rose, J. C. (2019). "Refining and Extending Task-Technology Fit Theory: Creation of Two Task-Technology Fit Scales and Empirical Clarification of the Construct," *Information & Management*, 56(5), pp.1-16.
- Jacoby, J. (2002). "Stimulus-Organism-Response

- Reconsidered: An Evolutionary Step in Modeling (Consumer) Behavior," *Journal of Consumer Psychology*, 12(1), pp.51-57.
- Jeyaraj, A. (2022). "A Meta-Regression of Task-Technology Fit in Information Systems Research," *International Journal of Information Management*, 65, pp.1-13.
- Jo, H. (2024). "Subscription Intention for ChatGPT Plus: A Look at User Satisfaction and Self-Efficacy," *Marketing Intelligence & Planning*, 42(6), pp.1052-1073.
- Kang, S., Choi, Y., and Kim, B. (2024). "Impact of Motivation Factors for Using Generative AI Services on Continuous Use Intention: Mediating Trust and Acceptance Attitude," *Social Sciences*, 13(475), pp.1-18.
- Lam, T. (2025). "Continuous Use of AI Technology: The Roles of Trust and Satisfaction," *Aslib Journal of Information Management*. Ahead-of-Print.
- Lee, J. C. and Chen, X. (2022). "Exploring Users' Adoption Intentions in the Evolution of Artificial Intelligence Mobile Banking Applications: The Intelligent and Anthropomorphic Perspectives," *International Journal of Banks Marketing*, 40(4), pp.631-658.
- Lee, C. T., Shen, Y. C., Wang, C. H., and Hung, H. Y. (2025). "Adapting to Generative AI: Examining the Users' Coping Strategies of Generative AI Image Systems," *Technological Forecasting & Social Change*, 218, pp.1-18.
- Malhotra, N. K., Kim, S. S., & Patil, A. (2006). "Common Method Variance in IS Research: A Comparison of Alternative Approaches and a Reanalysis of Past Research," *Management Science*, 52(12), 1865-1883.
- Mehrabian, A. and Russell, J. A. (1974). *An Approach to Environmental Psychology*, The MIT Press, Cambridge, MA.
- Moussawi, S., Koufaris, M., and Benbunan-Fich, R. (2021). "How Perceptions of Intelligence and Anthropomorphism Affect Adoption of Personal Intelligent Agents," *Electronic Markets*, 31, pp.343-364.
- Nunnally, J. C. (1967), *Psychometric Theory*, New York: McGraw-Hill.
- Podsakoff, P. M., MacKenzie, S. B., Lee, J-Y., & Podsakoff, N. P. (2003). "Common Method Biases in Behavioral Research: A Critical Review of the Literature and Recommended Remedies," *Journal of Applied Psychology*, 88(5), 879-903.
- Premathilake, G. W., Li, H., Li, C., Liu, Y., and Han, S. (2025). "Understanding the Effect of Anthropomorphic Features of Humanoid Social Robots on User Satisfaction: A Stimulus-Organism-Response Approach," *Industrial Management & Data Systems*, 125(2), pp. 768-796.
- Seifdar, M. H. and Amiri, B. (2025). "Strategic Adoption of Generative AI in Organizations: A Game-Theoretic and Network-Based Approach," *International Journal of Information Management*, 84, pp.1-19.
- Tiwari, C. K., Bhat, M. A., Khan, S. T., Subramaniam, R., and Khan, M. A. I. (2024). "What Drives Students toward ChatGPT? An Investigation of the Factors Influencing Adoption and Usage of ChatGPT," *Interactive Technology and Smart Education*, 21(3), pp.333-355.
- Truong, T. T. H. and Chen, J. S. (2025). "When Empathy is Enhanced by Human-AI Interaction: An Investigation of Anthropomorphism and Responsiveness on Customer Experience

- with AI Chatbots,” *Asia Pacific Journal of Marketing and Logistics*, Ahead-of-Print.
- Uddin, S. M. F., Kirmani, M. D., Bin Sabir, L., Faisal, M. N., and Rana, N. P. (2025). “Consumer Resistance to WhatsApp Payment System: Integrating Innovation Resistance Theory and SOR Framework,” *Marketing Intelligence & Planning*, 43(2), pp.393-411.
- Vafaei-Zadeh, A., Nikbin, D., Wong, S. L., and Hanifah, H. (2025). “Investigating Factors Influencing AI Customer Service Adoption: An Integrated Model of Stimulus-Organism-Response (SOR) and Task-Technology Fit (TTF) Theory,” *Asia Pacific Journal of Marketing and Logistics*, 37(6), pp.1465-1502.
- Verma, S., Kashive, N., and Gupta, A. (2025). “Examining Predictors of Generative-AI Acceptance and Usage in Academic Research: A Sequential Mixed-Methods Approach,” *Benchmarking: An International Journal*, Ahead-of-Print.
- Wang, H., Tao, D., Yu, N., and Qu, X. (2020). “Understanding Consumer Acceptance of Healthcare Wearable Devices: An Integrated Model of UTAUT and TTF,” *International Journal of Medical Informatics*, 139, pp.1-10.
- Woodworth, R. S. (1929). *Psychology*, Revised Edition. Henry Holt and Company, New York.
- Wu, B. and Chen, X. (2017). “Continuance Intention to Use MOOCs: Integrating the Technology Acceptance Model (TAM) and Task Technology fit (TTF) Model,” *Computers in Human Behavior*, 67, pp.221-232.
- Zhang, X., Qi, Z., Ma, L., and Zhang, G. (2025). “Assessing the Curvilinear Relationship in Employee Digital Performance: A Task-Technology Fit Perspective,” *International Journal of Human-Computer Interaction*, 41(4), pp.2615-2633.
- Zhou, T. and Ma, X. (2025). “Examining Generative AI User Continuance Intention based on the SOR Model,” *Aslib Journal of Information Management*, Ahead-of-Print.
- Zhou, S., Li, T., Yang, S., and Chen, Y. (2022). “What Drives Consumers’ Purchase Intention of Online Paid Knowledge? A Stimulus-Organism-Response Perspective,” *Electronic Commerce Research and Applications*, 52, pp.1-13.

-
- 저자 박현선은 현재 경북대학교 경영학부 BK교육연구단 계약교수로 재직중이며, 경북대학교 경영학부에서 경영학 박사학위를 취득하였다. 주요 관심분야는 디지털플랫폼, 정보보안, 클라우드, AI 등이 있다.
 - 저자 김상현은 미국 University of Mississippi에서 경영정보전공으로 박사학위를 취득하였으며, 현재 경북대학교 경영학부 교수로 재직 중이다. 주요 연구 분야는 정보보안, Human-Computer Interaction, AI 윤리, 클라우드 컴퓨팅 등이 있으며, 다수의 연구 실적이 International Journal of Information Management, Information and Management, Journal of Computer Information Systems.
 - 저자 이민영은 현재 경북대학교 경영학부에서 박사과정을 수료하였고 동 대학원에서 석사학위를 취득했다. 연구 관심 분야는 인공지능 서비스, 온라인 플랫폼, 지능 정보 시스템 및 Human-computer Interaction 등 이다.

〈부록 1〉 측정항목의 탐색적 요인분석 결과

문항	성분								
	1	2	3	4	5	6	7	8	9
ant1	0.795	0.116	0.034	0.076	0.072	0.097	0.073	0.069	0.062
ant2	0.762	0.138	0.066	0.108	0.093	0.065	0.077	0.056	0.076
ant3	0.818	0.111	0.061	0.103	0.066	0.069	0.063	0.082	0.071
per1	0.120	0.827	0.045	0.087	0.082	0.097	0.074	0.067	0.064
per2	0.138	0.810	0.039	0.081	0.070	0.086	0.075	0.071	0.066
per3	0.128	0.764	0.057	0.124	0.063	0.065	0.079	0.060	0.068
int1	0.036	0.021	0.802	0.153	0.033	0.032	0.015	0.018	0.023
int2	0.033	0.032	0.785	0.098	0.026	0.022	0.028	0.029	0.026
int3	0.217	0.034	0.814	0.069	0.035	0.020	0.018	0.016	0.021
rar1	0.024	0.040	0.068	0.814	0.123	0.034	0.023	0.031	0.033
rar2	0.025	0.033	0.076	0.807	0.105	0.022	0.026	0.018	0.037
rar3	0.035	0.043	0.062	0.779	0.093	0.019	0.021	0.030	0.032
acc1	0.046	0.041	0.044	0.068	0.869	0.036	0.023	0.018	0.026
acc2	0.041	0.028	0.040	0.052	0.856	0.031	0.025	0.021	0.027
acc3	0.044	0.039	0.032	0.060	0.810	0.035	0.024	0.020	0.022
sat2	0.043	0.036	0.020	0.034	0.037	0.895	0.057	0.021	0.018
sat3	0.035	0.027	0.025	0.047	0.032	0.906	0.042	0.019	0.016
sat4	0.048	0.032	0.026	0.036	0.040	0.880	0.046	0.024	0.027
tf1	0.037	0.038	0.033	0.035	0.044	0.028	0.814	0.019	0.022
tf2	0.032	0.024	0.030	0.021	0.018	0.031	0.878	0.025	0.021
tf3	0.034	0.040	0.027	0.022	0.019	0.026	0.905	0.027	0.025
tf4	0.038	0.026	0.018	0.029	0.022	0.032	0.869	0.040	0.031
tp1	0.127	0.114	0.123	0.093	0.080	0.081	0.104	0.867	0.034
tp2	0.112	0.117	0.118	0.080	0.086	0.082	0.107	0.903	0.040
tp3	0.113	0.112	0.124	0.091	0.087	0.089	0.113	0.872	0.038
tp4	0.109	0.110	0.122	0.079	0.093	0.086	0.117	0.854	0.043
icu1	0.117	0.114	0.126	0.076	0.090	0.087	0.111	0.032	0.886
icu2	0.115	0.123	0.129	0.095	0.085	0.091	0.104	0.025	0.899
icu3	0.109	0.110	0.123	0.085	0.084	0.093	0.108	0.023	0.910
icu4	0.114	0.118	0.125	0.083	0.080	0.090	0.113	0.028	0.891

주) 요인추출방법: 주성분 분석, sat1 항목은 최초 적합도 분석 결과 이후 삭제됨.

〈부록 2〉 Harman의 단일요인 검정

설명된 총분산

성분	초기 고유값			추출제곱합 적재값		
	전체	0.317	%누적	전체	%분산	%누적
1	12.144	31.409	31.409	12.144	31.409	31.409
2	2.737	13.389	44.798	2.737	13.389	44.798
3	2.323	9.192	53.990	2.323	9.192	53.990
4	1.783	6.753	60.743	1.783	6.753	60.743
5	1.630	5.761	66.504	1.630	5.761	66.504
6	1.434	4.625	71.129	1.434	4.625	71.129
7	1.229	3.964	75.093	1.229	3.964	75.093
8	1.061	3.423	78.516	1.061	3.423	78.516
9	1.018	3.282	81.798	1.018	3.282	81.798
10	0.792	2.556	84.354			
11	0.702	2.264	86.618			
12	0.573	2.004	88.622			
13	0.558	1.983	90.605			
14	0.492	1.756	92.361			
15	0.425	1.588	93.949			
16	0.382	1.233	95.182			
17	0.307	0.922	96.104			
18	0.262	0.770	96.874			
19	0.239	0.625	97.499			
20	0.159	0.511	98.010			
21	0.135	0.411	98.421			
22	0.127	0.366	98.787			
23	0.113	0.317	99.104			
24	0.095	0.287	99.391			
25	0.089	0.160	99.551			
26	0.049	0.144	99.695			
27	0.045	0.135	99.830			
28	0.042	0.073	99.903			
29	0.023	0.055	99.958			
30	0.017	0.042	100.000			