

ChatGPT 같은 멘토, 우리 같은 멘티: AI 멘토십의 특성과 영향

Mentor Like ChatGPT, Mentee Like Us: AI Mentorship's Characteristics and Influence

이원준(주저자)

Won-jun Lee(First Author)

청주대학교 경영학과 Cheongju University, Business Administration Department(marketing@cju.ac.kr)

인공지능(AI) 기술의 급속한 발전은 캐릭터 시뮬라크라를 통한 AI 기반 멘토링 또는 e-코칭을 기업이 활용할 수 있는 가능성을 열어주었다. 조직은 멘토십 프로그램 도입을 통하여 필요 기술의 개발을 촉진하고 구성원간 네트워킹 기회를 제공하며, 일과 삶의 균형에 대한 실질적 조언 제공이 가능할 것으로 기대된다. 그러나 이러한 이점에도 불구하고, 성공적인 멘토십 도입을 위해서는 자격을 갖춘 멘토의 확보, 멘토링에 수반되는 시간 제약과 같은 실행상의 문제들이 여전히 존재한다. 이러한 문제들을 해결하고 멘토십 프로그램의 효과성을 높이기 위하여 ChatGPT와 같은 대형 언어 모델 AI를 활용한 대안적 멘토링 전략이 주목받고 있다. 이 연구는 AI 기반 멘토링 시스템의 중요 속성을 확인하고, 이러한 속성들이 AI 멘토 효능감을 통해 멘티 만족도에 어떤 영향을 미치는지 실증적으로 분석하는 것을 목표로 한다. 또한 멘티 만족도에 미치는 의인화의 조절 효과의 확인을 통하여 멘토십 개발의 방향성을 제시하였다. 본 연구는 개념적 논의와 실증적 연구를 통해 AI 기반 멘토의 핵심 특성을 식별하고 향후 사용자 연구에 대한 필요한 실무적 시사점을 제공한다.

주제어: 인공지능, 대규모 언어모델, 캐릭터 시뮬라크라, e-코칭, 멘토링, 챗GPT

The rapid advancement of artificial intelligence (AI) technologies has introduced the potential for AI-based mentoring or e-coaching through the integration of character simulacra. Organizational mentorship programs are expected to enhance skill development, provide networking opportunities, and offer practical advice on work-life balance. However, despite these benefits, practical challenges persist, such as difficulty acquiring qualified mentors and the time constraints associated with mentoring. Consequently, an alternative mentorship strategy utilizing advanced AI based on large language models such as ChatGPT is gaining attention to address these challenges and enhance the effectiveness of mentorship programs. This study aims to identify the critical attributes of AI-based mentoring systems and empirically analyze how these attributes impact mentee satisfaction through AI mentor efficacy. Additionally, the moderating effect of anthropomorphism on satisfaction was examined. Through conceptual discussions and empirical research, this study identifies essential attribute variables of AI-based mentors and provides insights for future user research in mentoring services.

Keyword: AI, large language model, character simulacra, e-coaching, mentoring, ChatGPT

최초투고일: 2024. 05. 26

수정일: (1차: 2024. 07. 05)

게재확정일: 2024. 07. 08

Copyright 2024 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

비약적인 컴퓨팅 연산능력 발전과 인공지능의 급속한 부각은 직업 훈련이나 급여 관리, 교육 등 인력 관리의 제 분야에서의 업무 자동화(automation)의 가능성에 대한 관심을 증가시켰으며(Thakkar et al., 2020; Arora et al., 2023; Stojanov, 2023), 멘토(mentor)와 멘티(mentee)를 인적으로 연결해주는 멘토십(mentorship) 역시 AI 기술과 캐릭터 시뮬라크라(character simulacra)의 접목을 통하여 자동화가 가능한지에 대한 논의가 진전되고 있다(Khandelwal & Upadhyay, 2021; Abhari et al., 2023). 멘토링(mentoring) 혹은 코칭(coaching)은 전문적 지식을 갖춘 멘토와 멘티 간의 상호적인 관계 구축 활동을 의미하며(Nora & Crisp, 2007; Terblanche, 2020), 효과적인 멘토십 구축은 조직 인력의 전문성과 성장에 필수적인 요소이다. 비교적 경험이 풍부한 멘토와 경력이 짧은 멘티의 상호육성적인 관계를 조성함으로써 기업은 인재 양성의 효과성을 높일 수 있다.

다양한 이유에서 전문적 인적 자원의 육성 프로그램의 멘토십 구축 중요성이 강조된다. Bagai & Mane (2023)은 멘토십 제도 운영의 필요성으로 조직 내 업무와 개인 경력 발전에 필요한 기술 향상, 전문적인 인적자원 네트워킹의 기회 제공, 멘티에 대한 건설적 비판을 통한 자신감 강화, 가치있는 경력 기회를 제공하는 경력 상담을 제시하였다. 또한 멘토는 멘티에게 업무 뿐만 아니라 일과 삶의 균형(work-life balance)에 관한 효과적 조언을 제공할 수 있다고 주장하였다. 이와 같은 멘토링의 순기능적 영향으로 멘토링 프로그램 혹은 일대일 코칭을 지향하는 제도 및 프로그램은 기업, 대학 등 다양한 기관에 걸쳐 점

점 더 증가하고 있다(Köbis & Mehner, 2021).

그러나, 전통적 멘토십 프로그램이 강력한 이점을 제공하고 기업들도 필요성을 잘 인식하고 있음에도 불구하고 효과적인 멘토-멘티 관계의 형성은 어려운 과업이다. 일반적으로 우수한 멘토를 충분히 확보하거나 멘토링에 필요한 시간을 요청하기 어려운 경우가 많으며, 멘토-멘티 간 물리적 거리로 인한 지리적 장벽 존재, 일대일 맞춤 매칭의 어려움, 멘토십 제도 운영 후 성과 평가의 어려움 등이 존재한다(Cooke, 2018). 이에 우수한 멘토 확보의 어려움을 극복하고 멘토십의 효용성과 필요성을 확보하기 위하여 인간이 아닌 기술, 특히 인공지능(AI)을 멘토나 코치로 활용하는 대안적 멘토십 전략이 주목받기 시작하였다(Khandelwal & Upadhyay, 2021).

AI 멘토를 활용하는 대안적 접근법은 기술 기반 멘토십 구축 시 예상되는 한계와 장애물을 극복하기 위하여 머신러닝, 빅데이터 분석, 초거대 언어모델(LLM: large language model)과 같은 최첨단 AI 기술을 적극적으로 활용하고, 캐릭터 시뮬라크라 구축과 맞춤형 서비스를 추진한다. 캐릭터 시뮬라크라라는 인공적으로 제작된 인간과 유사한 캐릭터 혹은 존재를 의미하며(Baudrillard, 1994). 이제는 발전된 기술을 활용하여 인간의 외형뿐만 아니라 언어, 감정 등을 모방하고 인간처럼 상호작용하는 것이 가능해지고 있다. 이에 다양한 분야에서 AI가 멘토 역할을 수행하고 있으며, 예로 Rojas-Muñoz et al. (2021)은 멘토를 구하기 어려운 임상 의학 분야에서 이미 AI가 의료진을 대신하여 훌륭한 멘토 역할을 하고 있다고 주장하였다. AI 멘토십의 개발과 도입은 희소한 인간 전문가에 대한 과도한 의존을 줄이고 멘토의 본 업무에 지장을 초래하는 부담 요인을 경감함으로써 더 많은 기업과 조직들이 멘토십 제도를 확대하여 운영할 수 있게 해준다. 또한 AI 기반의 전문적

인 멘토링 플랫폼은 멘토-멘티 관계의 진척 상황을 휴먼 어널리틱스(human analytics) 도구를 활용하여 모니터링해주며, 동시에 성과 측정이 가능한 지표 생성을 통하여 멘토십 경험의 효과에 대한 통찰력을 제공할 수 있다.

그러나 이런 가능성에도 불구하고 이제 개념이 논의되고 구체화되기 시작한 초기의 혁신인 관계로 아직 AI 멘토십 프로그램의 경험자나 관련 학술적 연구는 희소하며, 특히 실증적 연구는 전무한 실정이다. 본 연구는 이러한 문제와 기회를 이해하고, AI 멘토의 도입이 어떻게 기존의 멘토링 방법론에 대한 대안을 제시할 수 있는지를 탐구함으로써, 현재 시점에서 중요한 연구주제임을 명확히 하고 AI 기반의 멘토십 도입과 운영이 조직내 인적 자원의 개발과 전문적 성장을 촉진할 수 있는지를 확인하고자 한다. 구체적으로는 실제 AI 멘토십 프로그램의 경험자들을 대상으로 AI 멘토십 프로그램의 성공 모델을 실증하고, 이를 통해 기업이나 조직에서 AI를 멘토리얼 적극 활용할 수 있는 방법을 제시하고자 한다.

II. 이론적 배경

2.1 초거대 인공지능과 AI 멘토

인공 지능(AI: artificial intelligence)은 인간의 인지적 기능을 대신 수행할 수 있는 컴퓨터 영상 기술, 자연어 처리, 로봇틱스, 자동화, 가상 현실 등 다양한 정보 기술들의 총체적 융합체다(Bughin & Hazan, 2017). 최근 진보된 AI 기술의 출현은 인간이 정보를 확보하고 문제를 해결하는 방식은 물론이고 인간-기계 간 상호작용 방식에도 일대 변혁을

초래하였다(Akiba & Fraboni, 2023). 한때 난제로 여겨졌던 튜링 테스트(Turing test)를 인간과 구분없는 대화로 통과하는 AI가 증가하였으며(Gonçalves, 2023), ChatGPT 같은 초대형 언어모델(LLM: large language model)도 상용화되기에 이르렀다(Stojanov, 2023). 특히 ChatGPT, LLaMA와 같은 초대형 언어모델(LLM)들은 강력한 자연어 처리 기술을 기반으로 가정과 직장, 일상 생활속에서 인간의 다양한 언어적 행동을 모방한 업무 수행의 가능성을 제시하였다(Park et al., 2023).

초대형 언어모델의 등장은 단순히 인간-기계 간의 질의 응답이나 정보 교환을 성공적으로 수행하는 것 이상의 활용 가능성을 보여주고 있다. Park et al. (2023)은 LLM(large language model)을 활용하여 기억을 합성하고 인간처럼 자연스럽게 행동하는 생성형 에이전트를 소개하였다. Bommasani et al. (2021)에 의하면 이런 에이전트들은 인간 행동에 대한 폭넓은 지식을 보유하고 있으며, 오랜 역사 동안 인간 사회에서 축적되어온 대규모의 사회 데이터로 훈련되었다. 기술은 또한 인간의 경험을 모방할 수 있는 ‘캐릭터 시뮬라크라(character simulacra)’의 가능성을 제시하였다. 캐릭터 시뮬라크라는 인간과 유사한 캐릭터나 인공적인 존재를 가리키는 용어이며, 이는 실제로 사람과 유사한 특성, 행동, 감정, 그리고 상호작용을 가진 인공적인 존재나 캐릭터를 구축하는 것을 의미한다(Baudrillard, 1994). 캐릭터 시뮬라크라는 사회 과학 연구(Riedl & Young, 2005), 게임 속 NPC 캐릭터 구축(Miyashita et al., 2017) 등 다양한 분야에서 인간과 비슷한 행동 및 상호작용을 모델링하거나 인간의 역할을 대신할 수 있는 기능을 제공할 수 있다(Madden & Logan, 2007; Shao et al., 2023).

시뮬라크라 연구는 음성 생성, 딥페이크(deep fake)

생성 등 다양한 방식으로 다중 유형(multi-modal)의 인간형 시뮬라크라 개발로 빠르게 진화하고 있다(Nguyen et al., 2022; Wang et al., 2022). 이와 관련하여 Shao et al.(2023)은 1) 캐릭터 프로필, 장면, 상호작용에 관한 경험 기반 데이터셋트 구축, 2) AI가 자신의 무지를 인정하고, 정보의 오류인 환각 현상(hallucination)을 보이지 않도록 예방하는 보호적 경험 구축, 3) 캐릭터별로 고유한 경험만 적용한 특정 경험의 업로드, 4) 캐릭터의 신뢰성과 타당성을 높이기 위한 기존 관행과의 비교라는 네 단계로 구성된 캐릭터 시뮬라크라 구축 방법론을 구체적으로 제시하였다. 이제 ChatGPT처럼 누구나 이용가능한 SaaS(software as a service)형 AI의 발전으로 캐릭터 시뮬라크라 개발과 이용이 더욱 손쉬워졌으며, 이를 통하여 기술이 인간과 유사한 감정을 느끼며 타인과의 상호작용을 기억하고 이에 맞춘 반응을 보이는 보다 생동감있는 역할 연기(role playing)가 가능해지고 있다. 이에 따라 인간이 아닌 컴퓨터 캐릭터 시뮬라크라로 구현된 멘토도 현실화되고 있으며, 궁극에는 인간 역할의 대체가 가능할 것으로 기대된다(Shao et al., 2023).

2.2 AI 멘토의 등장과 고려사항

최근 AI 기반 멘토링에 대한 기술적 가능성이 주목받고 있지만, 발전된 정보기술을 멘토링에 활용할 수 있다는 주장은 인터넷 등장 초기부터 제시되어왔다. 이와 관련하여 초기의 관련된 문헌 연구들은 e-코칭(e-coaching)이라는 용어를 사용해 왔다(Terblanche, 2020). Ribbers and Waringa(2015)는 e-코칭을 지리적으로 분리된 당사자들간에 비위계적 파트너십을 구축하는 것이며, 전통적 아날로그 수단과 가상 수단을 모두 활용하여 학습과 성찰 과정이 수행

되는 것으로 정의하였다. Clutterbuck(2010)은 e-mail, 기타 다양한 인터넷 자원들을 매개로 관계를 구축해나가는 과정을 e-코칭이라고 정의하였다. 이들의 초기 정의는 새롭게 활용되기 시작한 인터넷이나 전자 매체를 코칭이나 멘토링의 수단으로 주목하였다는 점에서 의의가 있다. 그러나, 기술을 단지 인간과 인간 사이를 중재하는 매개적 수단으로 한정적 이해하였다는 점, 커뮤니케이션 수단을 인터넷과 온라인으로 다양화하는 것에 국한하였다는 점 등의 한계도 존재한다. 그러나, e-코칭에 대한 관심은 이후 AI 기반 멘토링 분야의 진전에 새로운 가능성을 제시하게 되었다. 이에 AI 기술을 활용하여 멘토십이나 코칭 인프라를 구축하려는 노력은 최근 다양하게 시도되고 있다. 우선 먼저 손쉽게 접근할 수 있는 영역은 AI를 특정한 기술 영역에 한정한 후 전문적 멘토십을 구축하는 것이다. Rojas-Muñoz et al.(2021)은 종양 치료 분야에 AI 멘토 도입의 가능성을 제시하였으며, 수술 과정에 대한 다양한 사진 이미지들을 빅데이터로 구축하고 머신러닝으로 분석함에 따라 의료진이 참조할 수 있는 높은 정확도의 멘토링 시스템을 구축하였다. Bagai and Mane(2023)은 AI 멘토의 영역 확대를 주장하며 교육 기관, 기업 조직은 물론이고 전문가 협회, 정부 기관, 비영리 기관, 온라인 교육 플랫폼, 경력 코칭 서비스 그리고 온라인 원격 근무 환경 등 다양한 사례와 산업군에서 활용될 수 있을 것으로 전망하였다.

향후 AI 기반의 멘토십 활용 증대가 가져올 효과성과 위험 요인을 파악하기 위해서는 그 장단점을 확인할 필요가 있다. 우선 장점으로는 멘티 개개인에 대한 맞춤형 멘토링이 가능하고, 시간과 공간의 장벽없이 낮은 비용으로 신속하게 제공이 가능하다는 점이다. 보다 구체적으로 Bagai and Mane(2023)은 AI 멘토십 운영이 조직 구성원 개개인의 경력 성장을

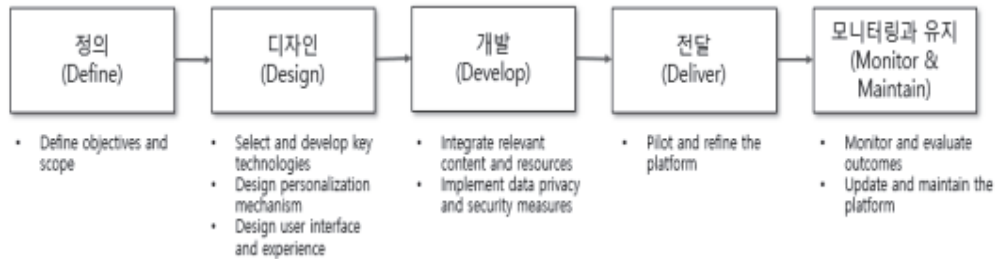
촉진하고, 새로운 업무 기술을 익히며, 일과 생활 간의 균형을 향상시켜 구성원 참여와 조직 만족도 향상에 긍정적 역할을 수행할 것으로 주장하였다. Arora et al.(2023)은 AI 멘토십의 주요 장점으로 확장성을 주장하였다. 전통적인 인적 기반 멘토링이 일대일 상호작용에 의존함에 따라 멘토가 효과적으로 지원할 수 있는 멘티의 수는 제한될 수 밖에 없으나 AI 멘토는 다수의 멘티를 동시다발적으로 지원할 수 있다. 이는 다양한 지역과 시간대에 걸쳐 멘토십에 대한 접근성을 강화하고, 정보의 포용성, 다양성을 촉진할 것으로 기대하였다. Arora et al.(2023)은 AI 멘토의 지속적인 학습과 적응 능력에 주목하였다. AI는 끊임없이 멘티의 진행 상황을 모니터링하고, 방대한 양의 데이터를 분석, 실시간 피드백을 제공함으로써 최신, 최적의 멘토십 경험을 제공한다. 이들은 또한 AI 멘토십이 다양한 학제간 협력과 지식 공유를 촉진할 가능성을 제시하였다. 멘티들은 다양한 데이터 소스와 학문 분야에서 아이디어를 도출할 수 있는 AI 멘토의 능력을 활용하며, 이를 통하여 다양한 시각을 얻고 복잡한 문제에 대한 혁신적, 창의적 해결방안을 제시할 것으로 기대하였다.

그러나 이런 가능성에도 불구하고 여전히 존재하는 기술적 도전과 윤리적인 고려사항은 장애요인이 될 수 있다. 개인 정보의 잠재적 보안 문제, 훈련 데이터의 특성상 특정 사용자 그룹에 대하여 불공평하거나 차별적 조치를 할 알고리즘 편향의 가능성, AI 멘토가 감성 지능을 확보하기 어려우며 인간 멘토 대체에 따른 윤리성 문제 등이 여러 연구자들에 의하여 단점으로 제기되기도 하였다. Köbis and Mehner (2021)는 멘토링 과정에서의 AI 사용은 단순히 관련 법과 규정(예: 데이터 보호 및 비차별성 관련 법규)을 준수해야 하는 것 뿐만 아니라 윤리적 기준 및 미해결된 문제(예: 데이터의 무결성, 시스템의 안전

및 보안, 기밀성, 편견 방지, 신뢰 보장 및 알고리즘의 투명성 등)에 대한 철저한 이해가 필요하다고 주장하였다. Moberg and Velazquez(2004)는 멘토가 멘티에 비하여 일반적으로 우월한 권력을 행사하기 때문에 멘티의 지식, 지혜 증진을 위한 발전적 지원 혜택을 제공하여야 하며, 구체적으로는 멘티에 대한 선의의 원칙, 권력 행사를 통한 가해의 금지, 멘티의 모든 행동에 동의를 구하고 자주성을 보장할 것, 차별을 피하고 멘티에게 잠재적 이익과 부담이 공정하게 분배되도록 할 것, 이해 관계의 충돌을 피하며, 멘티에게 애정을 주되 공정할 것 등이 요구된다고 주장하였다. AI가 우월한 지식과 영향력, 멘티에 대한 사전 정보를 가지고 멘토링에 참여하게 된다면 이런 요구사항과 의무들은 AI 멘토에게도 요구되는 것이 타당할 것이다.

2.3 AI 멘토의 특성과 구축 방안

AI 멘토십 플랫폼을 구축하고 개인화된 상담과 개입을 실현하기 위해서는 일련의 컴퓨터 기술과 IT 혁신의 접목이 필요하다. Bagai and Mane(2023)은 대화형 AI 인터페이스, 개인화 능력, AI와 머신러닝, 자연어 이해와 처리, 이용자 프로파일링과 데이터 분석, 외부 자원과의 통합성, 그리고 개인 프라이버시 보호와 보안 기술이 필요하다고 주장하였다. 이들의 주장에 의하면 AI 멘토의 목표와 서비스 범위를 설정하는 초기 작업에서 인사관리 부서나 조직 책임자의 참여, 머신러닝 기술, 컴퓨터 알고리즘의 이해, 사용자 니즈를 반영하는 개인화 메커니즘, 사용자 친화적 인터페이스 등이 포함되어야 한다. 또한 이후의 개발 과정에서는 멘토링에 필요한 기술 문서, 경력 관리와 관련된 논문과 자료, 삶과 일의 균형을 구현하는 도구 등이 통합되어야 하며 개인의 정보 프



〈Figure 1〉 AI 멘토 구축 과정 (Bagai and Mane 2023)

〈Table 1〉 챗봇 멘토의 설계 시 고려사항 (Terblanche 2020)

주요 속성	설계 시 고려사항
신뢰 (Trust)	'불쾌한 골짜기(uncanny valley)'의 발생 방지 보안과 데이터 프라이버시 보호의 우려의 제거 지속적인 챗봇의 성격(personality)을 유지
공감성 (Empathy)	인간적 이름을 사용하고, 인간다운 대화의 단서를 제공 다양한 분야에서 이용자의 선호와 기호 사항을 기억
투명성 (Transparency)	인간과 다른 우수한 측면을 부각 자기 공개(self-disclosure)를 실행 상담의 목적과 윤리적 기준을 명시
예측가능성 (Predictability)	지속적인 학습으로 인하여 기대되는 가능한 행동 변화를 기술 예측 가능한 성격과 인간다운 변동성(변덕) 간의 균형 찾기 상호작용에서 대화의 맥락을 이해하고 사용하기
신뢰성 (Reliability)	우아하게 실패할 줄 알기(모름을 인정) 챗봇의 성능과 신뢰성을 지속적 관찰 확증적인 메시지의 제공
역량 (Ability)	기존에 수립된 이론적 모델들의 활용 개인화를 사용하고 일반론적 답변을 하지 않기
선의 (Benevolence)	긍정적인 의도를 알리기 긍정적 태도와 분위기를 보여주기
통합성 (Integrity)	한계점에 대하여 분명하게 알리기 소개 메시지에 목적을 분명하게 제시하기

라이버시와 보안에 대한 우려가 해소되어야 한다. 마지막 단계인 전달과 모니터링 과정에서는 전사적 실행 이전에 충분한 파일럿 테스트와 개선 과정이 포함되어야 한다. 또한 실행 이후에도 지속적으로 성과를 모니터링하고 개선하는 지속적 작업이 필요하다.

기업이 조직 내 AI 멘토링 플랫폼 확보를 위한 전략으로 Terblanche(2020)는 AI 멘토 구축 프레임워크인 DAIC(designing AI coach) 방법론을 주장하였다. DAIC 프레임워크는 AI 기반 멘토링이나 코칭은 목표 달성, 웰빙, 자아 존중감 향상, 혹은 행

동 교정처럼 반드시 멘토링이 지향하는 특정 목표에 초점을 맞추어야 하며 이의 작동 기반이 되는 AI 챗봇이나 시스템은 검증된 이론적 모델에 기반할 것을 주장하였다. 또한 추가로 멘토링의 성공 여부를 예측할 판단할 수 있는 기능이나 준거가 멘토링 시스템과의 상호작용 모델에 포함되어 있어야 하며, 멘토링 시스템은 대화의 요구 사항이나 개입 방식에 있어서 누구나 수용할 수 있는 윤리적 지침에 위배되서는 안 된다고 주장하였다. Terblanche(2020)는 이를 위한 핵심 관계 속성으로 신뢰성, 공감성, 투명성, 예측 가능성, 신뢰성, 역량, 선의, 통합성의 요인들을 제시하였다.

III. 연구모형 및 가설 구축

3.1 연구의 프레임워크

본 연구는 Terblanche(2020)의 DAIC(Designing AI Coach) 이론에서 개념적으로 주장된 AI 멘토 구축 논의를 실증 연구를 위한 기본 프레임워크로 활용하였으며, 주요한 속성 요인들이 멘토십에 참가한 멘토의 만족도에 미치는 영향을 실증적으로 분석하고자 하였다. 모형에 포함될 속성 요인들은 동료 연구자 검토를 통하여 일차 검증되었으며, 이후 신뢰(trust), 신뢰성(reliability) 같이 상호 중복적 개념은 통합하고, 매개 변수와 실질적인 개념 차이가 없는 역량(ability) 변수를 제거하는 등의 조정을 통하여 최종적으로 다섯 개의 특성 요인들을 도출하였다.

또한 이들 다섯 개 속성 요인과 만족도간의 관계를 매개하는 변수로서 AI 멘토에 대한 효능감(AI mentor efficacy)을 사용하였다. 개인이 특정한 작

업을 수행할 능력에 대한 자신감이나 확신을 의미하는 일반적인 자기 효능감(self-efficacy)과 다르게(Bandura, 1994), 인간이 아닌 PC나 스마트폰 등의 정보기술이나 기기들이 특정한 작업을 수행할 능력에 대한 효능감인 컴퓨터 효능감(Compeau & Higgins, 1995; Howard et al., 2016), 휴대전화, 모바일 효능감(Park et al., 2014; Ali & Warraich, 2021), AI 효능감(Lee et al., 2023) 등이 존재한다. 이와 마찬가지로 AI 멘토에 대해서도 효능감을 지각하는 것이 가능하며 이를 매개로 AI 멘토십 경험에 대한 만족의 경로를 설명하고자 하였다. 마지막으로, AI 서비스 도입과 관련되어 다수 연구에서 관심 대상이었던 AI 의인화 수준의 조절적 효과를 검증하고자 하였다.

3.2 모형 및 가설 수립

신뢰는 불확실한 상호작용 환경 하에서 다른 당사자의 행동으로 얻을 수 있는 긍정적인 행동에 대한 기대이므로(Bhattacharya & Devinney, 1998), AI 멘토에 대한 신뢰는 멘토링을 통하여 문제를 해결할 수 있다는 기대인 효능감에 영향을 미칠 수 있을 것이다. 기존 연구에서는 의사에 대한 신뢰가 신체 및 정신 건강이나 삶의 질 개선과 긍정적 관계가 있다고 주장되었는데(Pre'au et al., 2004), 상담자를 다루는 점에서 의사와 유사성이 있는 멘토 역시 신뢰 형성을 통하여 상담자의 심리적 문제를 해결하고 효능감을 제공할 수 있을 것으로 기대된다. 보다 직접적으로 Ladegar and Gjerde(2014)은 리더십 코칭과 관련된 실증 연구를 통하여 조직내 하급자가 상급자에 대하여 지각하는 신뢰가 리더 역할 효능감(leader role efficacy)에 긍정적 영향을 미침을 확인하였다. 이런 연구결과들은 AI 멘토의 신뢰

성이 높을수록 멘티들은 AI가 자신의 문제를 해결할 수 있는 속성을 가지고 있다고 판단할 가능성이 높으며, 이에 다음과 같은 가설 1을 수립하였다.

가설 1: 신뢰성은 AI 멘토 효능감에 긍정적 영향을 미친다.

공감이 효능감 지각에 미치는 효과는 다양한 교육 환경에서 확인되었다. Bohecker and Horn(2016)은 대학생들을 대상으로 상담 서비스를 제공한 뒤, 학생과 상담자 간에 공감대가 형성된 경우에 더 높은 효능감이 발생함을 확인하였다. 또한 교육 분야에서 교사에 대한 학생들의 공감 정도와 교사 효능감 간에는 높은 상관관계가 있음이 확인되었다(Goroshit & Hen, 2016). Dionigi et al.(2020)은 이탈리아 병원의 심리치료 자원봉사자 집단들을 대상으로한 연구에서 타인 관점의 이해 및 공감 전략이 효능감에 긍정적 영향을 미침을 확인하였다. AI 역시 대규모 언어모델을 기반으로 학습이 이루어졌으며, 학습 데이터에는 대량의 감성 텍스트가 포함되어 있다. 따라서 AI 멘티도 공감에 필요한 분석 능력을 가지고 있으며, 이에 다음과 같은 가설 2를 수립하였다.

가설 2: 공감성은 AI 멘토 효능감에 긍정적 영향을 미친다.

자신에 대한 공개와 높은 윤리적 기준을 의미하는 투명성이 효능감에 미치는 영향은 Woodroof et al.(2020)에 의하여 확인되었다. 이들은 소셜 미디어 인플루언서가 자신에 대하여 투명하게 공개하고 소비자가 그 투명성을 인지하였을 때, 인플루언서의 판매 제품에 대한 효능감 기대와 구매 의도가 동시에 증가함을 확인하였다. Zia-ur-Rehman et al.(2021)

은 정보 공유를 통한 재무 정보의 투명성이 개인의 재무 관리 능력에 대한 효능감을 증가한다고 주장하였다. AI 멘토 역시 이용자들에게 자신의 멘토링 모델에 대한 소개 자료를 제공하고 있으며, 자신이 참조한 언어모델이나 직업들을 공개하고 있다. 이에 다음과 같은 가설 3을 수립하였다.

가설 3: 투명성은 AI 멘토 효능감에 긍정적 영향을 미친다.

AI 멘토의 긍정적 선의가 효능감에 미치는 영향에 대한 근거는 조직 리더십 연구에서 찾아볼 수 있다. Xia et al.(2022)은 선의에 기반한 리더십이 조직의 창의적 성과와 연결되어 있는 것은 물론이고 창의성에 대한 자기 효능감과도 긍정적 영향이 있음을 확인하였다. Xu et al.(2018)은 선의의 리더십이 구성원의 조직 몰입을 촉진할 뿐만 아니라 구성원이 자신의 역할과 책임을 수행할 능력에 대한 자신감인 역할-폭 자아 효능감(role-breadth self-efficacy)에도 긍정적 영향을 미침을 주장하였다. AI 멘토는 복합적 감정을 지니거나 인간적 감정 기록에 영향을 받는 인간 멘토와 달리 멘티의 문제 해결이라는 선의만을 목적으로 구성된 알고리즘으로 인식된다. 이에 이들 리더십 연구에서 제시된 선의와 효능감 간의 관계를 기반으로 다음과 같은 가설 4를 수립하였다.

가설 4: 선의는 AI 멘토 효능감에 긍정적 영향을 미친다.

Alifuddin and Widodo(2021)는 자신의 한계점과 목표를 명확하게 인지하는 능력인 통합성은 효능감과 더불어 만족에 유의한 영향을 미친다고 주장하였다. Erkutlu and Chafra(2020)는 리더가 보여

주는 통합성은 조직 구성원의 일탈적 행위를 감소시키고 보다 조직 시민적 행동을 촉진하며 도덕적 효능감을 증가시킴을 확인하였다. Hoy and Woolfolk (1993)는 건전하고 건강한 학교 분위기와 교육 기관의 정직성이 교수에 대한 효능감 지각에 긍정적 영향을 미침을 확인하였다. 기술의 산물인 AI 멘토는 선악이라는 개념이 존재하지 않으며 가치 중립적 일 것으로 기대된다. 이에 다음과 같은 가설 5를 수립하였다.

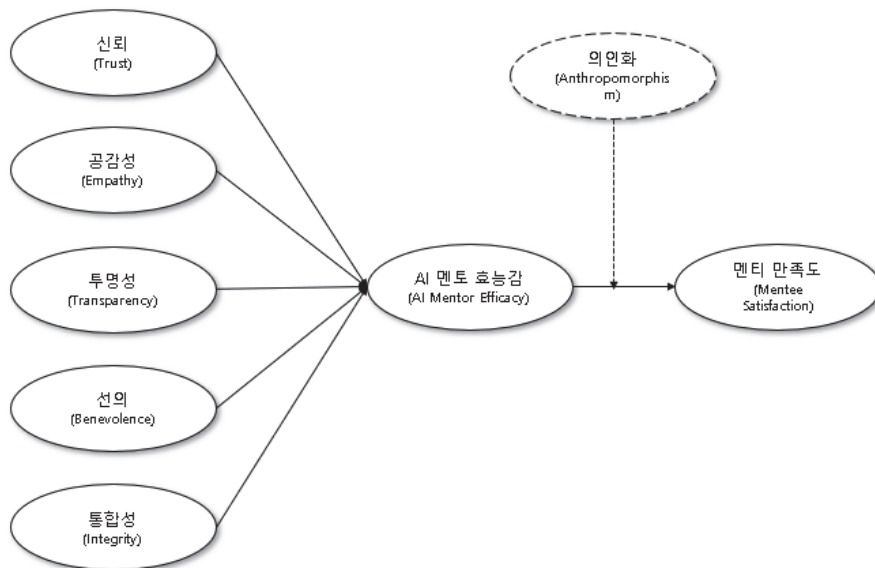
가설 5: 통합성은 AI 멘토 효능감에 긍정적 영향을 미친다.

고조된 효능감이 만족에 미치는 긍정적 영향은 다양한 환경에서 확인되었다. Moè et al.(2010)은 교사의 효능감이 직무 만족에 긍정적 영향을 미침을 확인하였고, Ellen et al.(1991)은 소비자의 효능감

이 구매 만족에 미치는 영향을 분석하였다. van de Ridder et al.(2015)은 의과대학 대학생들의 자기 효능감 지각이 업무 성과와 만족도에 긍정적 영향을 미침을 확인하였으며, 그 외 대학 교육, 직무 현장 등 많은 연구들에서 효능감과 만족도 간의 관계가 지지되었다(Doménech-Betoret et al., 2017; Bargsted et al., 2019). 특히 인간인 이용자는 오피스나 웹 브라우저 같은 단순한 프로그램의 우수한 성능에도 만족감을 느낄 수 있음을 고려할 때, AI 멘토의 효능감 지각이 만족에 영향을 미치리라 예상할 수 있다. 이에 다음과 같은 가설 6을 수립하였다.

가설 6: AI 멘토 효능감은 멘티 만족에 긍정적 영향을 미친다.

보통의 인간을 닮은 정도인 의인화의 조절 효과는 브랜드 캐릭터 연구(Zhang et al., 2020), 광고



〈Figure 2〉 연구 모형

(Reavey et al. 2018), 챗봇과 로봇(Blut et al., 2021). 컴퓨터 게임(Kim et al., 2016) 등 다양한 분야에서 검토되었으며, 이들 연구에서 의인화는 대부분 이용자의 태도를 강화시키는 정적 조절 효과를 가진 것으로 주장되었다. 예로, 명시적 대리인(EAK: explicit agent knowledge) 이론에 의하면 인간이 비인적 객체와 상호작용할 때 인간과 닮은 정도는 태도에 정적 영향을 미친다고 주장한다(Epley et al., 2007), 그러나, 이런 기존 연구들과 다르게 AI 멘토십의 맥락에서는 인간을 닮은 정도인 의인화가 반드시 긍정적 조절 효과를 나타낼 것이라고 확신하기는 어렵다. 불쾌한 골짜기(uncanny valley) 이론은 어설픈 인간 모방하기에 대하여 이용자들이 극심한 거부감을 느낄 수 있음을 주장한다(Wang et al., 2015). 또한 적지 않은 멘토링 분야의 연구들에서 멘티가 멘토에게 요구하는 모습은 문제 해결자로서 단순한 인간 이상의 비범한 능력을 기대하고 있음이 주장되고 있음을 고려해야 한다. 일 예로 Andersen and Reese(1999), Grassi(2013)는 멘토링은 삶의 방향을 제시하는 초월적인 영적 활동의 일환이라고 주장하였으며, Aymukhambet et al.(2020)은 멘토에게 기대되는 역할은 초자연적인 보호자이자 조언자 역할이라고 주장하였다. Robertson and Lawrence(2015)는 관계문화이론을 적용하여 멘토의 역할은 평범한 인간이 아니라 영웅의 여정에 비견된다고 주장하였다. 이상의 연구들은 적어도 멘토십 분야에서는 멘티는 멘토에게 평범 이상의 능력을 기대하며, 이에 따라 멘토가 보통 인간과 다를 바가 없을 수록 오히려 멘티에 대한 기대와 만족도가 낮아질 부적 관계의 가능성을 제시하고 있다. 멘토에 대한 멘티의 기대를 반영하여 다음과 같은 의인화의 조절효과 가설을 수립하였다.

조절효과 가설: 의인화는 AI 멘토 효능감과 멘티 만족간의 관계를 부적으로 조절할 것이다.

IV. 실증 분석

4.1 데이터 수집 및 분석

제시된 연구 모형과 가설의 검증을 위하여 20대 대학생 표본을 대상으로 설문이 시행되었다. AI 멘토링은 대학생 뿐만 아니라 직장인, 노령층 등 다양한 직종과 연령의 이용자가 기대되는 서비스지만, 다수 선행 연구들에서 20대 이용자들이 혁신적 서비스의 초기 이용자이자 구전 전파자로서의 중요성이 강조되는 점(Shankar & Narang, 2020), 교내 상담 기관을 통한 상담 경험이 풍부하다는 점(Murray et al., 2016), 그리고 이들이 향후 기업과 조직의 신규 구성원으로써 멘토링 노력이 집중되어야 한다는 점에서 표본 선정은 타당하다고 판단되었다. 아직 전반적인 기술 확산과 서비스 경험이 이루어지지 않은 연구 주제인 AI 멘토십에 대한 응답자들의 민감도를 높이기 위하여 연구 참여 대상자들은 설문 참여 전에 2주간 실제 AI 멘토를 선정하여 개인적 고민이나 경력 등 코칭이 필요한 주제들을 상담하도록 요청하였다. 이를 위하여 OpenAI의 생성형 AI 기술 기반의 멘토링이 가능한 서비스인 뤼튼(wrtn.ai)이 제공하는 심리상담사 '지우'를 이용하였으며, 각자 자신의 고민과 고충을 상담한 이후에 그 내역을 스크린 캡처하여 다이어리 형태로 수집하고 멘토링 활동의 증거로 제출하게 하였다.



〈Figure 3〉 AI 상담사와의 상담

이후 설문 조사와 데이터 수집은 2024년 3월 중 10일間に 걸쳐 온라인으로 진행되었으며, 총 311명의 참가자중 불성실 응답자를 제외하고 총 308명의 유효한 데이터가 수집되었다. 이들은 AI 상담 내역 다이어리 제출외에 추가로 상담한 AI 멘토에 대한 주요한 특징들을 묘사하도록 요구되었으며 이를 통과하지 못한 응답자는 불성실 응답자로 간주하였다. 최종적으로 채택된 응답자의 평균 연령은 만 22.6세였으며, 성별로는 남성(47.7%)과 여성(52.3%)으로 비교적 고르게 나타났다. 이들은 생성형 AI 서비스를 일평균 2.53회 이용하였으며, AI 멘토와의 상담을 통하여 당초 상담 목적을 달성한 정도는 100점 만점 기준으로 74.52점 수준이라고 응답하였다. 전반적으로 표본의 특성은 연구 모형 검증을 위한 대

상으로 적합하다고 판단되었다.

4.2 측정도구의 신뢰성과 타당성

연구 가설 검증의 사전 단계로 측정 문항의 신뢰성과 타당성 평가가 필요하다. 본 연구의 모든 척도와 문항들은 기존 선행 연구들을 참조하여 AI 멘토십 연구에 적합하도록 번역과 자구 수정을 거친 후 〈표 2〉와 같이 사용되었다. 모든 측정 문항들은 1점(전혀 아니다) ~ 5점(정말 그렇다)으로 구분된 리커트(Likert) 5점 척도를 통하여 측정되었다.

다항목으로 구성된 측정 문항들은 항목의 선별과 정교화 과정이 필요하다(Churchill, 1979). 이를 위하여 신뢰성 분석과 타당성 분석이 순차적으로 진행

〈Table 2〉 척도 및 측정 문항

구성 변수 (출처)	측정 문항
신뢰 (Agag & El-Masry, 2017)	AI 멘토는 신뢰할만하다 AI 멘토는 믿을만하다 AI 멘토는 상담자를 우선시할 것이라 믿는다
공감성 (Jones et al., 2022)	상담중 AI 멘토와 교감할 수 있다 상담중 AI 멘토와 함께 하는 느낌이다 상담중 AI 멘토와 생각과 느낌을 공유하였다
투명성 (Norman et al., 2010)	AI 멘토는 자신의 생각을 정확하게 표현하였다 AI 멘토는 자신이 의미하는 바를 정확하게 제시하였다 AI 멘토는 자신의 실수를 인정하고 개선하였다
선의 (Walabe, 2013)	AI 멘토는 나에게 많은 관심을 보였다 AI 멘토는 나의 중요 관심사를 잘 살펴준다 AI 멘토는 최선을 다해 나를 도와주려고 한다
통합성 (Walabe, 2013)	AI 멘토는 상담자를 공정하게 대해주었다 AI 멘토의 상담 내용은 명확하고 일관되었다 AI 멘토가 가진 상담 원칙은 바람직하다
AI 멘토 효능감 (Lee et al., 2023)	나는 AI 멘토를 통하여 문제를 잘 해결할 수 있을 것이다 나는 AI 멘토를 통하여 할 일을 잘 할 수 있을 것이다 나는 AI 멘토를 사용하여 상담 목표를 달성할 수 있을 것이다
멘티 만족도 (Lee et al., 2023)	나는 AI 멘토에 만족한다 나는 AI 멘토가 기대를 충족시켰다고 생각한다 나는 AI 멘토를 활용하는 것은 좋은 선택이라고 생각한다

〈Table 3〉 신뢰성 및 타당성

구성 변수	크론바하 알파	복합신뢰도	AVE
AI 멘토 효능감	0.875	0.923	0.800
선의	0.712	0.837	0.633
공감성	0.875	0.923	0.800
통합성	0.708	0.836	0.630
멘티 만족도	0.898	0.936	0.830
투명성	0.684	0.827	0.617
신뢰	0.790	0.880	0.714

되었다. 우선 신뢰성 분석은 크론바하 알파 계수, 복합 신뢰도, AVE(average variance extracted) 값을 통하여 진행되었으며, 크론바하 알파 및 복합

신뢰도는 0.7 이상, 그리고 AVE 값은 0.5이상으로 나타나 요구하는 적정 기준치를 모두 충족하는 것으로 나타났다(Churchill, 1979).

측정 문항의 타당성 분석을 위하여 확인적 요인분석(CFA: confirmatory factor analysis)을 실시하였다. 요인 분석은 탐색적 요인분석(EFA: exploratory factor analysis)이나 확인적 요인분석을 선택적으로 진행할 수 있는데, 본 연구에서 사용된 모든 변수와 측정 문항들은 이미 타당성과 신뢰성이 검증된 선행 연구로부터 차용되었기 때문에 확인적 요인분석이 보다 적합하다고 판단되었다(Suhr, 2006). 분석 결과 <표 4>에 의하면, 각 문항들의 표준 요인 부하량이 모두 유의하여 수렴 타당성이 입증되었다. 확인적 요인분석의 주요 적합도 지표 값은 $\chi^2=390.696$,

$d.f=168.000$, $p=0.000$, $CFI=0.943$, $NFI=0.905$, $GFI(AGFI)=0.889(0.847)$, $SRMR=0.060$, $RMSEA=0.066$ 으로 나타났으며, χ^2 의 p값 외에 전반적으로 수용가능한 수준으로 나타났다.

추가적 판별 타당성은 포넬-락커 분석으로 검증하였다. Fornell-Larcker(1981)는 판별 타당성을 확인하기 위하여 AVE의 제곱근 값을 사용할 수 있으며, 계산된 값이 다른 잠재 변인들의 값보다 더 클 때 판별 타당성이 확인된다고 주장하였다. 일 예로 <표 5>에서 공감성의 제곱근 값은 0.894이며, 이는 관련된 다른 잠재 변수들의 값(0.632, 0.514, 0.465,

<Table 4> 확인적 요인분석

측정 문항	표준화 계수	S.D	t-value	p-value
trust1	0.919	0.036	15.227	0.000**
trust2	0.902	0.032	20.27	0.000**
trust3	0.487	0.052	8.034	0.000**
empathy1	0.829	0.048	17.815	0.000**
empathy2	0.852	0.048	9.731	0.000**
empathy3	0.829	0.039	17.502	0.000**
transparency1	0.729	0.037	13.336	0.000**
transparency2	0.785	0.042	19.118	0.000**
transparency3	0.482	0.035	13.953	0.000**
benevolence1	0.604	0.047	17.095	0.000**
benevolence2	0.819	0.046	17.127	0.000**
benevolence3	0.629	0.039	19.683	0.000**
integrity1	0.564	0.049	10.585	0.000**
integrity2	0.761	0.034	19.731	0.000**
integrity3	0.704	0.049	8.732	0.000**
efficacy1	0.857	0.041	14.532	0.000**
efficacy2	0.834	0.038	11.186	0.000**
efficacy3	0.818	0.038	12.862	0.000**
satisfaction1	0.892	0.041	17.021	0.000**
satisfaction2	0.876	0.044	17.273	0.000**
satisfaction3	0.822	0.039	18.266	0.000**

〈Table 5〉 포널-락커 테스트

구성 변수	(1)	(2)	(3)	(4)	(5)	(6)	(7)
(1) AI 멘토 효능감	0.894						
(2) AI 멘토 효능감	0.497	0.796					
(3) 공감성	0.632	0.514	0.894				
(4) 통합성	0.408	0.614	0.465	0.794			
(5) 멘티 만족도	0.807	0.525	0.660	0.449	0.911		
(6) 투명성	0.507	0.543	0.457	0.575	0.538	0.785	
(7) 신뢰	0.630	0.512	0.620	0.531	0.655	0.517	0.845

0.660, 0.457, 0.620) 보다 크다. 다른 변수들에서도 동일한 결과를 관찰할 수 있어 판별 타당성은 확인될 수 있었다.

4.3 가설 검증

제시된 모형의 인과 관계를 분석하기 위하여 PLS (partial least square) 구조방정식 모형을 사용하였다. PLS는 선행 연구가 부족한 새로운 연구 주제이거나 충분한 표본 수집이 어려운 경우에 효과적인 인과관계 분석 방법이다(Lee et al., 2018). 연구 주제인 AI 멘토십은 선행 연구가 매우 희소하고, 실증이 거의 이루어지지 않았음을 고려할 때 PLS가 보다 적합한 분석 방법으로 판단되었다. SmartPLS s/w(version 4.0)를 통하여 총 500회의 부트스트

래핑 시뮬레이션을 반복 수행하여 가설을 검증하였다.

검증 결과 〈표 6〉에 의하면 총 6개 가설 중 5개 가설이 유의하게 채택되었으며, 가설 5는 기각되었다. 즉 신뢰, 공감성, 투명성, 선의가 AI 멘토 효능감에 유의한 영향을 미쳤으며, AI 멘토 효능감은 멘티의 만족도에 유의한 영향을 미쳤다. 그러나, 통합성은 AI 멘토 효능감에 유의한 영향을 미치지 못하였다. 이후 변수 간 인과 관계의 크기인 설명력 확인을 위하여 모형의 r^2 값을 확인하였다. 특성 변수들이 AI 멘토 효능감에 미치는 r^2 (조정된 r^2) 값은 0.524(0.516)로 나타났으며, AI 멘토 효능감이 멘티 만족도에 미치는 값은 0.651(0.649)로 확인되었다. 또한 AI 멘토 효능감의 매개적 효과는 기각된 변수인 통합성을 제외하고 다른 모든 특성 변수들과의 관계에서 유의한 것으로 〈표 7〉에서 확인되었다.

〈Table 6〉 가설 검증

가설	표준 계수	S.D	t-value	p-value
H1. 신뢰 → AI 멘토 효능감	0.318	0.066	4.839	0.000**
H2. 공감성 → AI 멘토 효능감	0.333	0.064	5.186	0.000**
H3. 투명성 → AI 멘토 효능감	0.177	0.062	2.85	0.004**
H4. 선의 → AI 멘토 효능감	0.124	0.054	2.295	0.022**
H5. 통합성 → AI 멘토 효능감	-0.093	0.055	1.679	0.093
H6. AI 멘토 효능감 → 멘티 만족도	0.807	0.027	29.955	0.000**

〈Table 7〉 매개효과

매개경로	표준 계수	S.D	t-value	p-value
선의 → AI 멘토 효능감 → 멘티 만족도	0.100	0.043	2.316	0.021**
공감성 → AI 멘토 효능감 → 멘티 만족도	0.269	0.054	4.945	0.000**
통합성 → AI 멘토 효능감 → 멘티 만족도	-0.075	0.045	1.675	0.094
투명성 → AI 멘토 효능감 → 멘티 만족도	0.143	0.050	2.841	0.005**
신뢰 → AI 멘토 효능감 → 멘티 만족도	0.257	0.055	4.664	0.000**

〈Table 8〉 의인화 조절효과 검증

조절 가설	표준 계수	S.D	t-value	p-value
의인화 → 멘티 만족도	0.242	0.039	6.150	0.000**
(의인화 * AI 멘토 효능감) → 멘티 만족도	-0.063	0.027	2.384	0.017**

추가로 AI 멘토 효능감과 멘티 만족도 간의 관계를 조절하는 변수로 AI 멘토의 의인화 정도를 투입하여 조절 효과 분석이 수행되었다. 〈표 8〉 분석 결과에 의하면 조절 변수의 부(-)적 효과는 예상한 바대로 유의하게 확인되었다

V. 결론 및 시사점

5.1 학문적 시사점

본 연구는 기업이나 기관의 멘토십 운영과 관련된 다양한 학문적, 실무적 시사점을 제공할 것으로 기대된다. 본 연구의 학문적 시사점은 다음과 같다. 첫째, 멘토링과 e-코칭의 연구 분야에서 급속하게 부상하는 AI의 영향력을 조망하고 새로운 연구 의제로 채택할 필요성을 제시하였다. 기존의 멘토십 연구들이 인적 멘토링에 기반하여 연구가 진행되어 왔으나, 최근 특정 영역에서 인간 이상의 능력을 보여주는 약

인공지능(weak AI)의 발전을 고려하지는 못하였다. e-멘토링 역시 인터넷과 온라인을 활용하는 비대면 인적 멘토링이라는 개념에 정체된 것에 비하여 발전된 생성형 AI의 등장은 인간 멘토 없이도 AI에 의한 멘토링 기술적으로 구현하고 있다. 이에 본 연구는 인적자원 개발에 필요한 AI 멘토의 가능성을 이론적, 실증적 접근을 통하여 조망하였다.

둘째, AI 멘토에 대한 개념적 논의와 더불어 그 속성에 대한 개념과 측정 방법을 제안하였다. Terblanche (2020)는 DAIC 프레임워크를 통하여 AI 코칭의 특성 요인으로 신뢰성, 공감성, 투명성, 예측가능성, 신뢰성, 역량, 선의, 통합성의 요인들을 제시하였다. 그러나 그의 연구는 개념적 논의에 그쳤으며, 구체적인 측정이나 멘토링 성과에 미치는 인과적 영향에 대한 논의는 전개하지 않았다. 본 연구는 관련 연구를 보다 심층적으로 전개하여 AI 멘토의 주요한 특성 요인들을 정제하고 구체적인 측정 도구들을 실증적으로 제시하였다. 이는 향후 다양한 AI 멘토 연구로의 확대에 필요한 기본적인 연구 기반을 제시할 것으로 기대된다.

셋째, AI 멘토 효능감의 개념을 제시하고, 특성 요인과 AI 멘토 효능감, 멘티 만족도 간의 인과적 관계를 규명하였다. 즉, AI 멘토에 대한 신뢰, 공감성, 투명성, 선의의 요인들은 AI 멘토 효능감에 긍정적 영향을 미치고, 이를 매개로 멘티 만족도에도 긍정적 영향을 미침을 확인하였다. 이를 통하여 AI를 통한 성공적인 멘토링에 필요한 이용자의 태도 형성 요인을 모형화하였다. 또한 기존의 효능감 관련 개념들이 자아 효능감에서 컴퓨터 효능감, 모바일 효능감 등으로 적용 대상이 확대되어 왔는데, 인간을 닮은 시뮬라크라를 구현하는 AI에 대한 효능감 지각 필요성을 제시하였다.

넷째, 다른 가설들이 유의하게 입증된 것과 다르게 통합성이 AI 멘토 효능감에 미치는 영향은 지지되지 않았다. 통합성이 AI 멘토가 공정하고 일관된 상담 원칙을 가지고 있는지를 측정하였음을 고려할 때, 몇가지 기각의 이유들이 추정될 수 있다. 우선 멘티들이 공정하고 일관된 상담보다는 자신의 상황과 문제에 적합한 다소 주관적이고 맞춤형 상담을 기대하였을 가능성이 있다. 만일 AI가 통합적인 상담 원칙만을 강조하거나 고수한다면 AI 멘토는 알고리즘에 따라 정형화된 단순한 프로그램으로 인식될 가능성이 있으며, 이는 멘토의 자질에 미달한다고 판단하였을 가능성이 있다. 또한 멘티들은 AI 멘토와의 상호작용에서 인간적 연결과 감정적 지원을 필요로 할 수 있으며, 통합성에 대한 강조는 이러한 측면에 대한 기대 충족을 어렵게 할 수 있다.

다섯째, AI 멘토링에 참여한 경험자들을 대상으로 실증 연구를 진행하였으며, 이는 AI의 짧은 상용화 역사와 실증의 어려움으로 인하여 아직까지 거의 이루어지지 않은 정량적 연구 접근법이었다. 실제 이용자 집단의 경험에 대한 측정은 AI 멘토링은 물론이고, 향후 다양한 AI 기반 서비스의 이용자 연구에 필요한

직관과 시사점을 제공할 것으로 기대한다.

마지막으로 여섯째, AI 멘토 효능감과 멘티 만족도 간의 사이에서 영향을 미치는 의인화의 부적 효과를 확인하였다. 일반적으로 게임 캐릭터나 광고 등의 분야에서 의인화가 이용자 만족에 미치는 효과는 인과적 관계 혹은 조절적 관계에서 대부분 정의 영향을 미치는 것으로 보고되어 왔다. 그러나 AI 멘토쉽에서 의인화 효과는 다른 연구 환경들과 다르게 부적 효과가 나타날 것으로 가정하였으며 유의하게 입증되었다. 이는 인간을 닮은 AI 구현에 대한 활발한 학문적 논의에도 불구하고 실제 AI 서비스의 이용자들은 인간보다 탁월한 경험을 제공할 수 있다면 굳이 인간의 형태에 집착하지 않고 있음을 보여준다. 이는 평범한 인간 이상의 존재에 의한 멘토링 요구로 이해될 수 있으며, 다른 챗봇이나 검색 서비스 등 기타 AI 응용 서비스 분야에서도 서비스 개발 방향에 대한 시사점을 제공할 수 있을 것으로 기대된다.

5.2 실무적 시사점

추가적으로 본 연구의 실무적 시사점은 다음과 같다. 첫째, AI를 기반으로 멘토링 역량을 강화할 수 있는 대안을 제시하였다. 기존의 멘토쉽 구축 노력은 우수 멘토라는 희소한 인적 자원에 기반할 수 밖에 없는 한계를 보였으며, 멘토의 자질에 따른 상이한 성과, 멘토링에 대한 수요와 공급의 불일치 등의 문제점이 발생하여왔다. 이에 대한 대안적 해결방안으로 본 연구에서는 AI 멘토를 도입함으로써 멘티가 만족할 수 있음을 실증을 통하여 제시하였다. 기업이나 조직은 AI 멘토의 도입을 통하여 인력 교육, 고충 상담, 지식과 기술 교육 등 다양한 분야에서 성과를 거둘 수 있을 것이며, 시간과 장소의 구애없는 상시적 멘토링은 인적 자원의 역량 향상과 복지 개선에도 활

용될 수 있을 것이다.

둘째, AI 멘토가 어떤 모습으로 구현될 것인가에 대한 시사점을 제공한다. 기존 챗봇 등 AI를 활용한 응용 서비스의 개발은 주로 처리 속도, 반응 시간, 반응성, 검색 정확도, 사용자 인터페이스 등 기술적 특성을 중심으로 전개되어왔다. 그러나 공감에 기반한 개인적이고 은밀한 상담이 필요한 멘토십 구축을 위해서는 기술적 요인 외에 멘티가 호응할 수 있는 감성적 요인의 중요성이 간과될 수 없다. 멘토에 대한 신뢰, 공감성, 선의 등 감성적 요인은 물론이고, 상담의 목적과 윤리적 기준을 설정하고 준수할 필요성 역시 크다. 본 연구에서는 AI 멘토의 주요한 특성으로 신뢰, 공감성, 투명성, 선의, 통합성의 요인을 제시한 바 있다. AI 멘토를 구축한 기업이나 기관에서는 실제 서비스를 제공하기 이전에 이들 요인을 측정할 수 있는 상담 질문을 직접 AI 멘토에게 제시하고, 답변의 적정성을 검토하는 검증 과정을 추가할 수 있다. 이를 통하여 AI 멘토가 주어진 역할을 수행할 수 있는지 판단할 수 있을 것이다.

셋째, AI 멘토 효능감과 멘티 만족도에 대한 실증적 인과 관계가 확인되었으며, 이는 단순히 기술적 호기심이 아니라 조직 성과 창출을 위한 도구로써 AI 멘토링이 이용가능함을 보여준다. 이미 현재 수준에서도 AI 기술의 완성도는 인간 멘티를 효율적으로 보완할 수 있는 수준에 올라있으며, 기업의 인적 자원 개발과 기업 경쟁력 강화를 위하여 AI 멘토십 프로그램의 도입이 필요함을 시사하였다. 기업과 조직 내 AI 멘토에 대한 수요를 확인하고, 구체적으로 어떤 분야와 업무에 적용할 수 있을지 검토하는 노력이 요구될 것이다.

넷째, AI 멘토의 구축이 단순히 기존 인간 멘토를 대신하는 수준에 머물 필요는 없으며, 새로운 가능성을 적극 개척해야 함을 확인하였다. 의인화의 부적

인 효과가 지지된 것은 문제 해결과 상담 목적 달성이 가능하다면 멘티들이 단순한 인간의 모방에는 큰 관심이 없음을 시사한다. 개인이 인생에서 갈구하는 미지의 해답을 찾기 위하여 종교나 초월자, 위인에 기대는 것처럼 AI 멘토에게도 인간 멘토 이상의 성과에 대한 기대를 품고 있을 수 있다. 단순한 데이터 분석과 정보 탐색을 넘어 지식과 지혜의 창출자로 미래의 AI 진화가 기대되는 만큼, AI 멘토십을 구축하는 기업들 역시 혁신적 기술 발전을 빠르게 반영하고 새로운 시각과 접근법을 통하여 인간 멘토의 역량을 초월하는 AI 멘토 구축에 노력하여야 할 것이다.

5.3 한계 및 향후 연구방향

본 연구의 함의와 시사점에도 불구하고, 일정한 제약 사항이 존재하며 향후 이의 극복을 위한 추가적 연구의 필요성이 제기될 수 있다. 첫째, 대학생 집단이 멘토링 참가 경험이 풍부하며 서비스 확산과 구전 활동 등 가치가 높은 집단이지만 향후 연구는 표본 대상의 확대가 필요할 수 있다. 대학생 집단과 더불어 직장인, 가정 주부, 노년층 등 일상에서 멘토링이 필요한 대상은 다양하기 때문에 표본 대상의 확대가 필요하며 이를 통하여 연구 결과의 일반화와 실무적 시사점을 확대할 수 있을 것이다.

둘째, 연구 프레임워크로 활용된 Terblanche(2020)의 DAIC(designing AI coach) 프레임워크에서는 구체적인 측정 도구들은 제시되지 않았다. 이에 실증 모델을 구축하기 위하여 다수의 관련 선행 연구들을 확인하고 이들 연구에서 사용된 측정 도구들을 도입하였다. 그러나 이들 선행 연구들은 연구 배경이나 연구 대상이 각기 상이하였다. 이를 보완하기 위하여 동료 연구자를 통한 점검, AI 상황에 적합하도록 자주 수정 등을 거쳤고 공통방법 편이(common

mehod bias) 예방을 위한 노력을 기울였다. 그럼에도 불구하고 일부 측정항목의 신뢰성은 재차 검증될 필요가 있으며, AI 멘토쉽의 특성을 측정할 수 있는 직접적 연구의 필요성도 제기된다. 이를 통하여 보다 안정적이고 신뢰성 높은 측정도구를 제공할 필요가 있으며, 향후 연구에서는 AI 멘토쉽의 측정 도구 개발을 위한 척도 연구가 필요할 것이다. 이에 향후에는 추가적 연구를 통하여 본 연구의 의의와 시사점이 확장될 수 있을 것으로 기대된다.

참고문헌

- Abhari, K., Bhullar, A., Le, J. and Sufi, N.(2023), "Advancing Employee Experience Management (EXM) Platforms," *Strategic HR Review*, 22 (3), pp.102-107.
- Agag, G. M. and El-Masry, A. A.(2017), "Why Do Consumers Trust Online Travel Websites? Drivers and Outcomes of Consumer Trust toward Online Travel Websites," *Journal of Travel Research*, 56(3), pp.347-369.
- Akiba, D. and Fraboni, M. C.(2023), "AI-supported Academic Advising: Exploring ChatGPT's Current State and Future Potential toward Student Empowerment," *Education Sciences*, 13(9), pp.885.
- Ali, I. and Warraich, N. F.(2021), "The Relationship between Mobile Self-efficacy and Mobile-based Personal Information Management Practices: A Systematic Review," *Library Hi Tech*, 39(1), pp.126-143.
- Alifuddin, M. and Widodo, W.(2021), "How Social Intelligence, Integrity, and Self-efficacy Affect Job Satisfaction: Empirical Evidence from Indonesia," *The Journal of Asian Finance, Economics and Business*, 8(7), pp.625-633.
- Anderson, K. R. and Reese, R. D.(1999), *Spiritual Mentoring: A Guide for Seeking Giving Direction*, Intervarsity Press.
- Arora, S., Tiwari, S., Negi, N., Pargaien, S. and Misra, A.(2023), "The Role of Artificial Intelligence in Mentoring Students," in 2023 1st International Conference on Circuits, Power and Intelligent Systems (CCPIS), pp.1-6.
- Aymukhambet, Z. A., Mirzakhmetov, A. A., Aituganova, S., Mukhazhanova, R. M. and Karipbaev, Z.(2020), "The Primodal Patron or Mentor, Archetypes of the Mythological World," *PalArch's Journal of Archaeology of Egypt/Egyptology*, 17(6), pp.13111-13120.
- Bagai, R. and Mane, V.(2023), "Designing an AI-Powered Mentorship Platform for Professional Development: Opportunities and Challenges," *International Journal of Computer Trends and Technology*, 71(4), pp.108-114.
- Bandura, A.(1994), *Self-efficacy*, Encyclopedia of human Behavior.
- Bargsted, M., Ramírez-Vielma, R. and Yeves, J. (2019), "Professional Self-efficacy and Job Satisfaction: The Mediator Role of Work Design," *Revista de Psicología del Trabajo y de las Organizaciones*, 35(3), pp.157-163.
- Baudrillard, J.(1994), *Simulacra and Simulation*, University of Michigan Press.
- Bhattacharya, R., Devinney, T. M. and Pillutla, M. M.(1998), "A Formal Model of Trust based on Outcomes," *Academy of Management Review*, 23(3), pp.459-472.
- Blut, M., Wang, C., Wunderlich, N. V. and Brock,

- C.(2021), "Understanding Anthropomorphism in Service Provision: A Meta-analysis of Physical Robots, Chatbots, and Other AI," *Journal of the Academy of Marketing Science*, 49, pp.632-658.
- Bohecker, L. and Doughty Horn, E. A.(2016), "Increasing Students' Empathy and Counseling Self-efficacy through a Mindfulness Experiential Small Group," *The Journal for Specialists in Group Work*, 41(4), pp.312-333.
- Bommasani, R., Hudson, D. A., Adeli, E., Altman, R., Arora, S., von Arx, S., ... and Liang, P. (2021), "On the Opportunities and Risks of Foundation Models," arXiv preprint arXiv: 2108.07258.
- Bughin, J. and Hazan, E.(2017), "The New Spring of Artificial Intelligence: A Few Early Economies," *VoxEU*. org, 21.
- Churchill, G. A. Jr.(1979), "A Paradigm for Developing Better Measure of Marketing Constructs," *Journal of Marketing Research*, 16(1), pp. 64-73.
- Clutterbuck, D.(2010), "Coaching Reflection: The Liberated Coach," *Coaching: An International Journal of Theory, Research and Practice*, 3(1), pp.73-81.
- Cooke, D. M.(2018), *The Influences of Professional Development, in Educative Mentoring, on Mentors' Learning and Mentoring Practices*, Doctoral dissertation, Auckland University of Technology.
- Compeau, D. R. and Higgins, C. A.(1995), "Computer Self-efficacy: Development of a Measure and Initial Test," *MIS Quarterly*, 19(2), pp.189-211.
- Dionigi, A., Casu, G. and Gremigni, P.(2020), "Associations of Self-efficacy, Optimism, and Empathy with Psychological Health in Healthcare Volunteers," *International Journal of Environmental Research and Public Health*, 17(16), 6001.
- Doménech-Betoret, F., Abellán-Roselló, L. and Gómez-Artiga, A.(2017), "Self-efficacy, Satisfaction, and Academic Achievement: The Mediator Role of Students' Expectancy-value Beliefs," *Frontiers in Psychology*, 8, 277668.
- Ellen, P. S., Bearden, W. O. and Sharma, S.(1991), "Resistance to Technological Innovations: An Examination of the Role of Self-efficacy and Performance Satisfaction," *Journal of the Academy of Marketing Science*, 19, pp. 297-307.
- Epley, N., Waytz, A. and Cacioppo, J. T.(2007), "On Seeing Human: A Three-factor Theory of Anthropomorphism," *Psychological Review*, 114(4), 864.
- Erkutlu, H. and Chafra, J.(2020), "Leader's Integrity and Interpersonal Deviance: The Mediating Role of Moral Efficacy and the Moderating Role of Moral Identity," *International Journal of Emerging Markets*, 15(3), pp.611-627.
- Fornell, C. and Larcker, D. F.(1981), "Evaluating Structural Equation Models with Unobservable Variables and Measurement Error," *Journal of Marketing Research*, 18(1), 39-50.
- Gonçalves, B.(2023), "Can Machines Think? The Controversy That Led to the Turing Test," *AI & SOCIETY*, 38(6), pp.2499-2509.
- Goroshit, M. and Hen, M.(2016), "Teachers' Empathy: Can it be Predicted by Self-efficacy?," *Teachers and Teaching*, 22(7), pp.805-818.
- Grassi, J.(2013). *The Spiritual Mentor*, Thomas Nelson.
- Howard, S. K., Ma, J. and Yang, J.(2016), "Student

- Rules: Exploring Patterns of Students' Computer-efficacy and Engagement with Digital Technologies in Learning," *Computers & Education*, 101, pp.29-42.
- Hoy, W. K. and Woolfolk, A. E.(1993), "Teachers' Sense of Efficacy and the Organizational Health of Schools," *The Elementary School Journal*, 93(4), pp.355-372.
- Jones, C. L. E., Hancock, T., Kazandjian, B. and Voorhees, C. M.(2022), "Engaging the Avatar: The Effects of Authenticity Signals during Chat-based Service Recoveries," *Journal of Business Research*, 144, pp.703-716.
- Khandelwal, K. and Upadhyay, A. K.(2021), "The Advent of Artificial Intelligence-based Coaching," *Strategic HR Review*, 20(4), pp.137-140.
- Kim, S., Chen, R. P. and Zhang, K.(2016), "Anthropomorphized Helpers Undermine Autonomy and Enjoyment in Computer Games," *Journal of Consumer Research*, 43(2), pp.282-302.
- Köbis, L. and Mehner, C.(2021), "Ethical Questions Raised by AI-supported Mentoring in Higher Education," *Frontiers in Artificial Intelligence*, 4, 624050.
- Ladegard, G. and Gjerde, S.(2014), "Leadership Coaching, Leader Role-efficacy, and Trust in Subordinates. A Mixed Methods Study Assessing Leadership Coaching as a Leadership Development Tool," *The Leadership Quarterly*, 25(4), pp.631-646.
- Lee, C. H., Chiang, H. S., and Hsiao, K. L.(2018), "What Drives Stickiness in Location-based AR Games? An Examination of Flow and Satisfaction," *Telematics and Informatics*, 35(7), pp.1958-1970.
- Lee, W. J., Lee, H. S. and Cha, M. K.(2023), "AI Like ChatGPT, Users Like Us: How ChatGPT Drivers and AI Efficacy Affect Consumer Behaviour," *Virtual Economics*, 6(4), pp. 44-59.
- Madden, N. and Logan, B.(2007), Collaborative Narrative Generation in Persistent Virtual Environments. In AAAI Fall Symposium: Intelligent Narrative Technologies, pp.71-78.
- Miyashita, S., Lian, X., Zeng, X., Matsubara, T. and Uehara, K.(2017), Developing Game AI Agent Behaving Like Human by Mixing Reinforcement Learning and Supervised Learning, In 2017 18th IEEE/ACIS International Conference on Software Engineering, Artificial Intelligence, Networking and Parallel/Distributed Computing (SNPD), pp. 489-494.
- Moberg, D. J. and Velasquez, M.(2004), "The Ethics of Mentoring," *Business Ethics Quarterly*, 14(1), pp.95-122.
- Moè, A., Pazzaglia, F. and Ronconi, L.(2010), "When Being Able is not Enough. The Combined Value of Positive Affect and Self-efficacy for Job Satisfaction in Teaching," *Teaching and Teacher Education*, 26(5), pp.1145-1153.
- Murray, A. L., McKenzie, K., Murray, K. R. and Richelieu, M.(2016), "An Analysis of the Effectiveness of University Counselling Services," *British Journal of Guidance & Counselling*, 44(1), pp.130-139.
- Nguyen, T. T., Nguyen, Q. V. H., Nguyen, D. T., Nguyen, D. T., Huynh-The, T., Nahavandi, S., ... and Nguyen, C. M.(2022), "Deep Learning for Deepfakes Creation and Detection: A Survey," *Computer Vision and Image Understanding*, 223, 103525.
- Nora, A. and Crisp, G.(2007), "Mentoring Students:

- Conceptualizing and Validating the Multi-dimensions of a Support System,” *Journal of College Student Retention: Research, Theory & Practice*, 9(3), pp.337-356.
- Norman, S. M., Avolio, B. J. and Luthans, F.(2010), “The Impact of Positivity and Transparency on Trust in Leaders and Their Perceived Effectiveness,” *The Leadership Quarterly*, 21(3), pp.350-364.
- Park, J. S., O’Brien, J., Cai, C. J., Morris, M. R., Liang, P. and Bernstein, M. S.(2023), Generative Agents: Interactive Simulacra of Human Behavior, In Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology, pp.1-22.
- Park, L. G., Howie Esquivel, J. and Dracup, K. (2014), “A Quantitative Systematic Review of the Efficacy of Mobile Phone Interventions to Improve Medication Adherence,” *Journal of Advanced Nursing*, 70(9), pp.1932-1953.
- Pre’au, M., Leport, C., Salmon-Ceron, D., Carrieri, P., Portier, H., Chene, G., et al.(2004), “Health-related Quality of Life and Patient-provider Relationships in HIV infected Patients during the First Three Years after Starting PI-containing Aantiretroviral Treatment,” *AIDS Care*, 16(5), pp.649 - 661.
- Reavey, B., Puzakova, M., Larsen Andras, T. and Kwak, H.(2018), “The Multidimensionality of Anthropomorphism in Advertising: The Moderating Roles of Cognitive Busyness and Assertive Language,” *International Journal of Advertising*, 37(3), pp.440-462.
- Ribbers, A. and Waringa, A.(2015), *E-Coaching: Theory and Practice for a New Online Approach to Coaching*, Routledge.
- Riedl, M. O. and Young, R. M.(2005), An Objective Character Believability Evaluation Procedure for Multi-agent Story Generation Systems, In International Workshop on Intelligent Virtual Agents (pp.278-291). Berlin, Heidelberg: Springer Berlin Heidelberg.
- Robertson, D. L. and Lawrence, C.(2015), “Heroes and Mentors: A Consideration of Relational-cultural Theory and “The Hero’s Journey,”” *Journal of Creativity in Mental Health*, 10 (3), pp.264-277.
- Rojas-Muñoz, E., Couperus, K. and Wachs, J. P. (2021), “The AI-Medic: An Artificial Intelligent Mentor for Trauma Surgery,” *Computer Methods in Biomechanics and Biomedical Engineering: Imaging & Visualization*, 9(3), pp.313-321.
- Shao, Y., Li, L., Dai, J. and Qiu, X.(2023), “Character-LLM: A Trainable Agent for Role-Playing,” *arXiv preprint arXiv:2310.10158*.
- Stojanov, A.(2023), “Learning with ChatGPT 3.5 as a More Knowledgeable Other: An Autoethnographic Study,” *International Journal of Educational Technology in Higher Education*, 20(1), pp.20-35.
- Shankar, V. and Narang, U.(2020), “Emerging Market Innovations: Unique and Differential Drivers, Practitioner Implications, and Research Agenda,” *Journal of the Academy of Marketing Science*, 48, pp.1030-1052.
- Suhr, D. D.(2006), “Exploratory or Confirmatory Factor Analysis?,” *Statistics and Data Analysis*, Paper 200-31, pp.1-17.
- Terblanche, N.(2020), “A Design Framework to Create Artificial Intelligence Coaches,” *International Journal of Evidence Based Coaching & Mentoring*,18(2), pp.152-162.
- Thakkar, D., Kumar, N. and Sambasivan, N.(2020),

- Towards an AI-powered Future That Works for Vocational Workers, In Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems, pp.1-13.
- van de Ridder, J. M., Peters, C. M., Stokking, K. M., de Ru, J. A. and Ten Cate, O. T. J.(2015), "Framing of Feedback Impacts Student's Satisfaction, Self-efficacy and Performance," *Advances in Health Sciences Education*, 20, pp.803-816.
- Walabe, E.(2013), *Trust in e-mentoring Relationships*, University of Ottawa (Canada).
- Wang, S., Lilienfeld, S. O. and Rochat, P.(2015), "The Uncanny Valley: Existence and Explanations," *Review of General Psychology*, 19(4), pp.393-407.
- Wang, Y., Kordi, Y., Mishra, S., Liu, A., Smith, N. A., Khashabi, D. and Hajishirzi, H.(2022), "Self-instruct: Aligning Language Model with Self Generated Instructions," *arXiv preprint*, arXiv:2212.10560.
- Woodroof, P. J., Howie, K. M., Syrdal, H. A. and VanMeter, R.(2020), "What's Done in the Dark Will be brought to the Light: Effects of Influencer Transparency on Product Efficacy and Purchase Intentions," *Journal of Product & Brand Management*, 29(5), pp.675-688.
- Xia, Z., Yu, H. and Yang, F.(2022), "Benevolent Leadership and Team Creative Performance: Creative Self-efficacy and Ppenness to Experience," *Frontiers in Psychology*, 12, 745991.
- Xu, Q., Zhao, Y., Xi, M. and Zhao, S.(2018), "Impact of Benevolent Leadership on Follower Taking Charge: Roles of Work Engagement and Role-breadth Self-efficacy," *Chinese Management Studies*, 12(4), pp.741-755.
- Zhang, M., Li, L., Ye, Y., Qin, K. and Zhong, J.(2020), "The Effect of Brand Anthropomorphism, Brand Distinctiveness, and Warmth on Brand Attitude: A Mediated Moderation Model," *Journal of Consumer Behaviour*, 19(5), pp. 523-536.
- Zia-ur-Rehman, M., Latif, K., Mohsin, M., Hussain, Z., Baig, S. A. and Imtiaz, I.(2021), "How Perceived Information Transparency and Psychological Attitude Impact on the Financial Well-being: Mediating Role of Financial Self-efficacy," *Business Process Management Journal*, 27(6), pp.1836-1853.

-
- 저자 이원준은 현재 청주대학교 경영학과 마케팅 전공 교수로 재직 중이다. 서울대학교 경영대학에서 마케팅 전공으로 경영학 박사학위를 받았다. Asis Pacific Journal of Marketing and Logistics, Virtual Economics, Journal of Theoretical and Applied Electronic Commerce Research, 마케팅 연구 등에 논문을 게재하였다. 정보기술에 대한 관심을 배경으로 온라인 소비자 행동, 디지털 마케팅, 구전 마케팅 등 다양한 분야의 연구를 진행하여 왔으며, 최근에는 인공지능 시대의 마케팅과 변화된 소비자 행동에 관하여 논문을 게재하였다.