

생성형 AI 리터러시가 학습 웰빙에 미치는 영향:

결과 의존도와 비판적 상호작용의 매개 및 윤리적 민감성의 조절효과

김 명 한*

동국대학교 행정대학원 대우교수

Impact of Generative AI Literacy on Learning Well-Being: The Mediating Roles of Result Dependence and Critical Interaction, and the Moderating Role of Ethical Sensitivity

Myoung-Han Kim*

Professor, Graduate School of Public Administration, Dongguk University, Seoul 04620, Korea

[요 약]

본 연구는 고등교육 현장에서 생성형 AI의 보편화에 주목하여, 대학생 330명을 대상으로 생성형 AI 리터러시가 학습 웰빙에 미치는 구조적 영향력을 규명하고자 조절된 이중매개모형을 검증하였다. 구조방정식 모델 분석 결과, 생성형 AI 리터러시는 학습 웰빙에 직접적인 영향을 미치기보다 결과 의존도를 유의하게 낮추고 비판적 상호작용을 촉진함으로써 학습 웰빙을 높이는 완전매개 효과를 나타냈다. 특히 학습자의 윤리적 민감성은 이러한 매개 경로를 유의미하게 조절하는 역할을 수행하였다. 윤리적 민감성이 높은 학생일수록 AI에 맹목적으로 의존할 때 '윤리적 부조화'(맹목적 의존에 따른 웰빙 저하)를 경험하는 반면, 비판적 상호작용을 수행할 때는 '윤리적 시너지'를 창출하여 웰빙 증진 효과가 극대화되는 양상을 보였다. 본 연구의 결과는 AI시대의 학습웰빙을 실현하기 위해서는 단순한 기술적 숙련도를 넘어 비판적 성찰 역량과 내면화된 윤리 의식을 체계적으로 결합한 통합적 교육 전략이 고등교육 차원에서 필수적으로 수립되어야 함을 시사한다.

[Abstract]

This study examines the effects of generative AI literacy on the learning well-being of 330 university students using a moderated dual-mediation model. Structural equation modeling revealed that generative AI literacy indirectly enhanced learning well-being by reducing result dependence and promoting critical interaction, indicating a full mediation effect. Notably, ethical sensitivity served as a significant moderator. Students with high ethical sensitivity experienced "ethical dissonance" (diminished learning well-being from excessive reliance) when exhibiting high result dependence, while achieving "ethical synergy" through critical interaction, which maximized their learning well-being. These findings suggest that fostering learning well-being in the AI era requires an integrated educational strategy. Beyond technical proficiency, higher education must harmonize critical interaction and internalized ethical sensitivity to effectively support students navigating this rapidly evolving technological landscape.

색인어 : 생성형 AI 리터러시, 학습 웰빙, 결과 의존도, 비판적 상호작용, 윤리적 민감성

Keyword : Generative AI Literacy, Learning Well-Being, Result Dependence, Critical Interaction, Ethical Sensitivity

<http://dx.doi.org/10.9728/dcs.2026.27.6.1595>



This is an Open Access article distributed under the terms of the Creative Commons Attribution Non-Commercial License (<http://creativecommons.org/licenses/by-nc/3.0/>) which permits unrestricted non-commercial use, distribution, and reproduction in any medium, provided the original work is properly cited.

Received 05 April 2026; Revised 11 May 2026

Accepted 29 May 2026

*First Author; Myoung-Han Kim

Tel: 

E-mail: okmhkim@gmail.com

1. 서 론

1-1 연구의 배경 및 필요성

2022년 말 대화형 인공지능(ChatGPT)의 등장을 기점으로 촉발된 생성형 인공지능(Generative AI) 기술의 비약적 발전은 고등교육 생태계의 패러다임을 근본적으로 변화시키고 있다[1],[2]. 텍스트, 코드, 이미지를 아우르는 멀티모달(Multimodal) AI가 보편화됨에 따라, 대학생들은 정보 탐색을 넘어 보고서 작성, 데이터 분석, 아이디어 기획 등 인지적 작업의 전 과정에서 생성형 AI를 핵심적인 학업 도구로 수용하고 있다[3],[4].

초기 기술 수용 단계에서 학계의 주된 관심은 생성형 AI가 제공하는 도구적 효용성과 학업 성취도 향상이라는 기능적 측면, 혹은 튜터링 프로그램으로서의 가능성에 집중되었다[5],[6]. 그러나 생성형 AI가 학생들의 일상적 인지 과정을 점진적으로 대체하는 현시점에 이르러, 이러한 기술적 밀착이 학습자의 고차원적 사고 능력과 ‘학습 웰빙(Learning Well-being)’에 어떠한 구조적 영향을 미치고 있는지에 대한 생태학적이고 심층적인 진단이 시급히 요구되고 있다. 최근 다수의 문헌은 학습자가 시스템의 한계를 비판적으로 검토하지 못한 채 생성형 AI에 의존할 경우 발생하는 ‘인지적 오프로딩(Cognitive Offloading)’ 현상의 위험성을 경고한다[7]. 즉, 대학생들이 복잡한 과제를 스스로 해결하려는 인지적 노력을 회피하고 AI의 산출물에 전적으로 의존하는 ‘결과 의존도(Result Dependency)’를 보일 경우, 단기적인 효율성은 얻을 수 있으나 장기적으로는 학업적 자기효능감의 저하와 통제권 상실로 인한 학업 스트레스의 가중을 경험하게 된다는 것이다[8]-[10].

이와 대조적으로, 생성형 AI를 맹목적인 정답지가 아닌 아이디어를 고도화하고 자신의 논리적 허점을 점검하는 ‘비판적 대화 파트너’로 활용할 경우, 학습자의 학업적 궤적은 전혀 다른 양상을 띤다[11],[12]. AI 시스템이 생성한 정보의 환각 현상(Hallucination)과 편향성을 인지하고 이를 주도적으로 수정해 나가는 ‘비판적 상호작용(Critical Interaction)’은 학생들의 메타인지와 문제해결 능력을 향상시키며[13],[14], 나아가 학습 성과와 학습 웰빙을 유의하게 증진시키는 긍정적 효과가 실증된 바 있다[15],[16].

이러한 두 가지 상반된 활용 궤적(결과 의존도 vs. 비판적 상호작용)을 결정짓는 주요 선행 변인으로 최근 학계는 학습자의 주도적 통제 역량인 ‘AI 리터러시(AI Literacy)’에 주목하고 있다[17]-[19]. 본 연구가 제기하는 가장 핵심적인 학술적 문제의식은, AI 리터러시에 따른 활용 방식의 차이가 학습 웰빙에 미치는 영향이 단순히 인지적 차원에 머물지 않으며, 학습자 개인에 내재된 ‘윤리적 민감성(Ethical Sensitivity)’ 수준에 따라 복잡한 심리적 역설을 발생시킨다는 점이다.

고등교육 현장에서 생성형 AI의 무분별한 사용은 진정성 있는 평가(Authentic Assessment)의 신뢰성을 위협하고 학업적 진실성(Academic Integrity)을 훼손하는 심각한 윤리적 딜레마를 초래하고 있다[20],[21]. 성적 압박과 과제 부담에 직면한 학생들의 AI 수용 의도 이면에는 복잡한 윤리적, 행동적 요인들이 교차하고 있다[22]. 이러한 맥락을 심리학의 고전적 패러다임인 ‘인지적 부조화 이론(Cognitive Dissonance Theory)’[23]에 적용해 보면 새로운 인과적 추론이 성립한다.

학업적 정직성에 대한 윤리적 민감성이 높은 학생일수록, 불가피한 외적 요인으로 인해 생성형 AI에 맹목적으로 의존하게 되었을 때 자신의 확고한 도덕적 신념과 실제 행동 간의 심리적 불일치를 경험하게 된다. 즉, 학자적 양심을 저버렸다는 자괴감과 결합하여 학업 스트레스가 유의미하게 가중되는 심리적 부작용(본 연구에서는 이를 ‘윤리적 부조화’ 현상으로 명명함)을 겪게 될 것으로 예상할 수 있다. 반면, 동일하게 높은 윤리적 민감성을 지닌 학생이 AI 산출물의 오류를 주도적으로 통제하는 비판적 상호작용을 수행했을 경우, 학업적 무결성(Integrity)을 달성했다는 내적 정당성을 획득하여 학습 웰빙이 증진되는 효과(본 연구에서는 이를 ‘윤리적 시너지’ 현상으로 명명함)가 나타날 것으로 기대된다.

그럼에도 불구하고 기존의 다수 문헌들은 생성형 AI에 대한 의존을 ‘규범적 이탈’이라는 징벌적 관점(Punitive approach)에서 주로 접근해 왔으며[24], 윤리 의식이 투철한 학생들이 활용 방식에 따라 겪는 심리적 역동의 구조를 정량적으로 밝혀낸 실증 연구는 부족한 실정이다.

따라서 본 연구는 전국 대학생 표본을 대상으로 생성형 AI 환경에서의 AI 리터러시가 활용 방식(결과 의존도, 비판적 상호작용)을 매개하여 학습 웰빙에 미치는 직·간접 경로를 분석하고, 이 구조적 관계 속에서 학습자의 ‘윤리적 민감성’이 조절변수로서 어떻게 기능하는지를 조절된 이중 매개모형(Dual Moderated Mediation Model)을 통해 실증적으로 검증하고자 한다. 본 연구의 결과는 고등교육 현장에서 양심적인 학습자들을 보호하고 능동적 협력을 유도할 수 있는 AI 윤리 교육 가이드라인 및 평가 방식 개편을 위한 정책적 시사점을 제공할 것이다.

1-2 연구의 목적 및 연구 문제

본 연구의 궁극적인 목적은 대학생의 AI 리터러시가 ‘결과 의존도’와 ‘비판적 상호작용’이라는 상이한 활용 방식을 매개하여 ‘학습 웰빙’에 미치는 총체적(직접 및 간접) 영향을 분석하고, 이 과정에서 학습자의 ‘윤리적 민감성’이 조절변수로서 어떠한 심리적 상호작용 효과를 일으키는 지 구조방정식 모형(SEM)을 통해 규명하는 데 있다. 본 연구가 수행하고자 하는 구체적인 연구 문제는 다음과 같다.

연구 문제 1: 대학생의 AI 리터러시는 생성형 AI에 대한

‘결과 의존도’와 ‘비판적 상호작용’에 각각 어떠한 영향을 미치는가?

연구 문제 2: 생성형 AI에 대한 ‘결과 의존도’와 ‘비판적 상호작용’은 학습자의 ‘학습 웰빙’에 각각 어떠한 영향을 미치는가?

연구 문제 3: 대학생의 AI 리터러시는 학습 웰빙에 직접적인 영향을 미치는가, 아니면 활용 방식(결과 의존도, 비판적 상호작용)을 통한 간접효과(매개효과)만을 가지는가?

연구 문제 4: 학습자의 ‘윤리적 민감성’은 생성형 AI 활용 방식(결과 의존도, 비판적 상호작용)과 학습 웰빙 간의 관계에서 유의미한 이중 조절효과를 나타내는가?

II. 이론적 배경 및 가설 설정

2-1 생성형 AI 리터러시의 진화와 인지적 오프로딩

1) 생성형 AI 리터러시와 할루시네이션(Hallucination)에 대한 비판적 인식

최근 생성형 AI가 학생들의 고차원적 인지 과정을 직접적으로 대행하기 시작하면서, 리터러시의 개념은 단순한 도구 활용을 넘어 비판적 사고(Critical Thinking)와 인지적 개입(Cognitive Engagement)을 포괄하는 다차원적 역량으로 진화하였다[15],[17]. 특히 학계는 AI 시스템이 허위 정보를 생성해 내는 ‘할루시네이션(Hallucination)’ 현상에 주목하며, 이러한 오류가 정보 수용 체계에 인지적 왜곡을 초래할 수 있음을 경고한다[25]–[27]. 따라서 현재 고등교육 맥락에서의 AI 리터러시는 생성형 AI가 산출한 정보의 편향성을 인지하고 지식 산출의 주도권을 통제하는 방어 기제 역량으로 규정된다[18],[28].

2) 결과 의존도 vs. 비판적 상호작용

본 연구에서 ‘결과 의존도’란 학습자가 인지적 부하를 회피하기 위해 생성형 AI가 산출한 결과물을 비판적 여과 없이 맹목적으로 수용하고 학업에 그대로 적용하려는 성향으로 정의된다[7]. 반면, ‘비판적 상호작용’은 AI를 절대적 권위자가 아닌 협력적 도구로 인식하고, 산출된 정보의 편향성을 능동적으로 검증하며 자신의 사고를 확장해 나가는 주도적 개입 과정으로 정의된다[11].

이러한 AI 리터러시 수준은 활용 방식을 두 가지 상반된 경로로 분기시킨다. 리터러시 역량이 부족한 학습자는 AI 산출물의 논리적 허점을 판별할 능력이 부재하기 때문에, 자신의 인지적 노력을 AI 시스템에 전적으로 위임하려는 ‘결과 의존도(Result Dependency)’를 띠게 되며, 이는 무비판적인 인지적 오프로딩(Cognitive Offloading)의 함정으로 이어진다[7],[8]. 반면, 시스템의 한계를 명확히 인지하는 리터러시를 갖춘 학습자는 생성형 AI에 맹목적으로 의존하지 않는다.

이들은 AI를 아이디어 점검을 위한 대화 파트너로 전유(Appropriation)하여, AI 시스템과의 역할 분담을 능동적으로 조율하는 ‘비판적 상호작용(Critical Interaction)’을 수행한다[12],[29].

가설 1a: 대학생의 AI 리터러시는 생성형 AI에 대한 ‘결과 의존도’에 유의미한 부(-)의 영향을 미칠 것이다.

가설 1b: 대학생의 AI 리터러시는 생성형 AI와의 ‘비판적 상호작용’에 유의미한 정(+)의 영향을 미칠 것이다.

2-2 비판적 상호작용과 학습 웰빙의 다차원적 구조

1) 결과 의존도의 역설과 학업적 불안 가중

무비판적인 ‘결과 의존도’는 장기적으로 학업적 부작용을 초래한다. 인지적 오프로딩이 심화될수록 학습자의 문제해결 능력은 퇴보하며 탈숙련화(Deskilling)에 대한 불안이 발생한다[30]. 내재화 과정을 생략한 학습자는 평가 환경에서 결과물의 논리를 방어하지 못할 것이라는 압박감에 시달리게 되며[31], 이러한 통제권 상실은 학업 스트레스를 가중시켜 학습 웰빙의 유의미한 저하를 초래한다.

가설 2a: 생성형 AI에 대한 ‘결과 의존도’는 학습 웰빙에 유의미한 부(-)의 영향을 미칠 것이다.

2) 비판적 상호작용의 심리적 보상

이와 대조적으로, 생성형 AI와 ‘비판적 상호작용’을 수행하는 학습자는 긍정적인 심리적 쾌감을 경험한다. 초기 인지적 과부하를 상호작용을 통해 경감시키고 논리를 재구성하는 과정은 메타인지와 자기조절 효능감을 향상시킨다[13],[32]. 이러한 통제 경험은 막막함을 해소하고 학업적 성취감과 웰빙을 증진시킨다.

가설 2b: 생성형 AI와의 ‘비판적 상호작용’은 학습 웰빙에 유의미한 정(+)의 영향을 미칠 것이다.

3) AI 리터러시와 학습 웰빙의 인과적 간극과 이중 매개 가설

일반적으로 기술 수용에 대한 통제력인 AI 리터러시 역량을 갖춘 학습자는 신기술 도입에 따른 불안감을 경감시킬 수 있으므로, 학습 웰빙에 직접적인 정(+)의 영향을 미칠 것으로 예상할 수 있다(가설 3a). 그러나 본 연구는 더 나아가 AI 리터러시가 학습 웰빙으로 직결되는 경로 이면에 존재하는 인과적 간극(Causal Gap)에 주목한다. AI 시스템의 원리와 한계에 대한 ‘인지적 역량’이 올바른 ‘행동(Behavior)’으로 발현되지 않는다면 그 심리적 보상은 제한적일 수 있다. 즉, 지식이 풍부하더라도 맹목적인 ‘결과 의존도’를 보인다면 웰빙은 훼손될 것이며, AI 시스템과 통제권을 나누는 ‘비판적 상호작용’이라는 실천적 행동을 거칠 때 비로소 웰빙이 극대화될 것이다. 따라서 본 연구는 AI 리터러시가 학습 웰빙에 미치는 영향이 활용 방식을 거쳐 전달되는 이중 매개 구조를 가질 것으로 추론한다.

가설 3a: 대학생의 AI 리터러시는 학습 웰빙에 정(+)의 직접적인 영향을 미칠 것이다.

가설 3b: AI 리터러시는 ‘결과 의존도’를 매개로 학습 웰빙에 부(-)의 간접 영향을 미칠 것이다.

가설 3c: AI 리터러시는 ‘비판적 상호작용’을 매개로 학습 웰빙에 정(+)의 간접 영향을 미칠 것이다.

2-3 윤리적 민감성의 조절효과: 부조화 및 시너지 현상의 추론

본 연구에서 ‘윤리적 민감성’은 생성형 AI 활용 과정에서 발생할 수 있는 표절, 정보 왜곡 등의 윤리적 딜레마를 인지하고, 도덕적 주체로서 책임감을 느끼는 심리적·인지적 반응 능력을 의미한다. 치열한 과제 부담에 직면한 대학생들의 기술 수용 의도 이면에는 이러한 복잡한 윤리적 딜레마가 교차하고 있다[22],[24],[25].

특히, 최근 인공지능 윤리 관련 선행연구들은 학습자의 윤리적 민감성이 앞서 논의한 두 인지적 태도(결과 의존도와 비판적 상호작용)의 발현 양상을 결정짓는 핵심 기제임을 지지한다. 이중과정 이론(Dual-Process Theory)의 관점에서 볼 때, 윤리적 책임감이 높은 학습자일수록 단순한 인지적 오프로딩에 해당하는 ‘결과 의존도’에 대해 강한 심리적 저항감을 느끼며 이를 억제하는 부적(-) 상관관계를 보인다[8],[25]. 반면, 윤리적 민감성은 정보의 출처를 확인하고 환각(Hallucination) 오류를 능동적으로 교정하려는 ‘비판적 상호작용’과는 강한 정(+)의 상관관계를 갖는다[12],[29],[34]. 즉, 학습자의 윤리적 민감성은 활용 방식이 학습 웰빙으로 이어지는 인과 경로에서 매우 중요한 심리적 통제(조절) 기제로 작동하며, 이는 다음과 같은 두 가지 대비되는 조절 효과(부조화 및 시너지)로 추론될 수 있다.

1) 결과 의존도의 조절효과: ‘윤리적 부조화(Ethical Dissonance)’ 현상 예상

Festinger의 인지적 부조화 이론에 따르면, 개인은 자신의 확고한 내적 신념과 상충하는 행동을 수행할 때 심리적 스트레스를 경험한다[23]. 무단 도용을 경계하고 평가 공정성(윤리적 민감성)을 중시하는 학생일수록, 불가피하게 생성형 AI에 의존하게 되었을 때 자신의 신념과 위배되는 행동을 취했다는 사실로 인해 유의미한 죄책감을 느끼게 된다[25]. 즉, 높은 윤리적 민감성은 AI 의존이 초래하는 웰빙의 저하를 더욱 악화시키는 조절변수로 작용할 것이며, 본 연구는 이를 ‘윤리적 부조화’ 현상으로 명명한다.

2) 비판적 상호작용의 조절효과: ‘윤리적 시너지(Ethical Synergy)’ 현상 예상

반면, 높은 윤리적 민감성이 비판적 상호작용과 결합할 때는 정반대의 상호작용 효과가 예상된다. 윤리 의식이 철저한 학생이 AI 산출물을 수동적으로 모방하지 않고 인간의 비판적 사유[33]를 발휘하여 AI 시스템을 통제해 냈을 때, 이들

은 자신의 학자적 양심을 수호해 냈다는 내적 정당성을 획득하게 된다[34]. 진실성을 지켜낸 긍정적인 협업 경험은 학습자의 웰빙을 더욱 가속화시키는 조절 변수로 작용할 것이며, 본 연구는 이를 ‘윤리적 시너지’ 현상으로 기대한다.

가설 4a: 학습자의 ‘윤리적 민감성’은 결과 의존도가 학습 웰빙에 미치는 부(-)의 영향을 더욱 가중시킬 것이다 (윤리적 부조화 현상 예상).

가설 4b: 학습자의 ‘윤리적 민감성’은 비판적 상호작용이 학습 웰빙에 미치는 정(+)의 영향을 더욱 증진시킬 것이다 (윤리적 시너지 현상 예상).

III. 연구 설계

3-1 연구 모형

본 연구는 대학생의 AI 리터러시가 상반된 활용 방식을 매개하여 학습 웰빙에 미치는 직·간접 영향을 분석하고, 이 과정에서 조절변수인 ‘윤리적 민감성’이 유발하는 이중 조절효과를 실증적으로 규명하기 위해 다음과 같이 조절된 이중 매개모형(Dual Moderated Mediation Model)을 구축하였다.

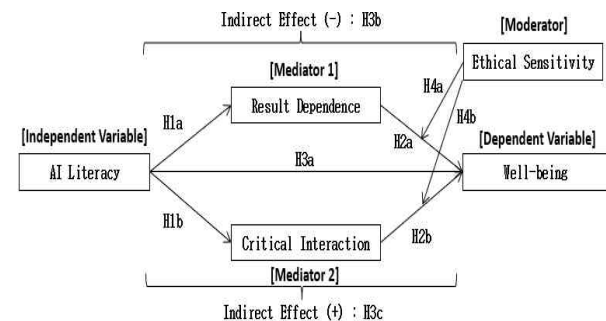


그림 1. 조절된 이중 매개모형

Fig. 1. Moderated dual mediation model (research model)

3-2 변수의 조작적 정의 및 측정도구

본 연구에 사용된 잠재변수의 측정도구는 기존 문헌에서 타당성이 검증된 척도를 고등교육 환경의 생성형 AI 활용 맥락에 맞게 수정 및 보완하여 사용하였다. 모든 문항은 리커트(Likert) 5점 척도(1: 전혀 그렇지 않다 ~ 5: 매우 그렇다)로 측정되었으며, 구체적인 조작적 정의와 측정 문항의 구성은 다음과 같다.

첫째, 본 연구에서 생성형 AI 리터러시는 기기를 조작하는 단순한 기술적 능력을 넘어, AI 산출물의 한계와 편향성을 인지하고 이를 비판적으로 평가하며 학업 목적에 맞게 주도적으로 통제할 수 있는 다차원적 인지 역량으로 정의한다. 이를 측정하기 위해 [17],[18],[19]의 문항을 본 연구에 맞게 수정하여 총 5개 문항으로 구성하였다.

둘째, 결과 의존도는 학습자가 인지적 노력을 회피(Cognitive offloading)하고, AI가 도출한 결과물을 무비판적으로 수용하여 학업 수행의 주도권을 AI에게 양도하려는 맹목적 의존 성향으로 정의한다. 이는 [8],[9],[10]을 바탕으로 2개 문항으로 구성하였다.

셋째, 비판적 상호작용은 AI의 산출물을 최종 정답이 아닌 ‘초안’이나 ‘참고 자료’로 인식하고, 학습자 스스로 팩트 체크, 논리적 검증, 정보의 재구성을 수행하는 능동적이고 주도적인 학습 과정으로 정의하며, [12],[14],[29]의 문항을 수정하여 2개 문항을 사용하였다.

특히, 위 두 변수(‘결과 의존도’, ‘비판적 상호작용’)를 각각 2개 문항으로 구성한 것은 해당 개념의 단일차원성(Unidimensionality)이 매우 강하기 때문이다. 무리하게 문항을 추가할 경우 발생할 수 있는 의미 중복(Semantic redundancy)과 측정 오차를 방지하고 내용 타당도를 극대화하기 위해 핵심 지표만을 엄선하였다[35]. 실증 분석 결과에서도 요인적재량과 신뢰도(CR, AVE)가 기준치를 상회하여 척도의 적합성이 검증되었다.

넷째, 학습 웰빙은 단순한 학업적 성취나 일시적 만족감을 넘어, 학업 수행 과정에서 경험하는 높은 자기효능감, 긍정적인 정서, 그리고 지적 성장이 동반되는 최적의 심리적 충만(Flourishing) 상태로 정의하며, [15],[16],[32]의 척도를 보완하여 3개 문항으로 구성하였다.

다섯째, ‘윤리적 민감성’은 AI 활용 시 발생할 수 있는 표절, 정보 왜곡 등의 윤리적 문제를 인지하고 도덕적 주체로서 책임감을 느끼는 심리적 상태로 정의하고, [20],[24],[33]의 척도를 수정하여 5개 문항으로 측정하였다.

3-3 자료 수집 및 분석 방법

본 연구의 자료 수집은 2025년 11월부터 2026년 1월까지 국내 수도권 및 비수도권에 소재한 7개 대학 소속 교·강사진

의 연구 협조를 통해 진행되었다. 구체적인 모집 경로(Recruitment path)는 해당 대학의 강의를 수강하는 학생들에게 온라인 설문조사 접속 링크를 배포하여 자발적인 참여를 유도하는 방식으로 이루어졌다. 연구의 타당성을 확보하기 위해 표본의 포함 기준(Inclusion criteria)은 생성형 AI를 학업에 활용하는 ‘정규 학위과정의 학부 재학생’으로 한정하였다. 반면, 학업적 맥락이 다른 대학원생 및 최고위과정 등 비학위과정 수강생은 분석 대상에서 제외(Exclusion criteria)하였다. 아울러 데이터 정제 과정에서 모든 문항에 동일한 척도로 응답한 획일적 답변(Straight-lining), 유사 혹은 상반된 문항 간에 논리적으로 모순되는 극단치 답변 등 14명의 불성실 응답을 추가로 배제하고, 최종적으로 330명의 유효 표본만을 실증 분석에 활용하였다.

수집된 자료는 SPSS 27.0과 AMOS 27.0을 활용하여 인구통계학적 특성에 대한 빈도분석 및 일원배치 분산분석(ANOVA)을 실시하였으며, 확인적 요인분석(CFA), 신뢰도 및 타당도 검증, 구조방정식 모형(SEM)을 통한 경로 분석, 그리고 부트스트래핑(Bootstrapping, N=5,000)을 적용한 간접효과 검증을 수행하였다. 조절효과 분석 시 다중공선성을 통제하기 위해 모든 변수는 파이썬(Python)을 활용하여 상호작용항(Interaction Term)으로 생성하여 평균 중심화(Mean-centering) 처리를 거쳤다.

3-4 조사 대상자의 일반적 특성

설문에 참여한 330명의 인구통계학적 특성은 표 2와 같다. 성별은 남성 168명(50.9%), 여성 162명(49.1%)이었으며, 전공은 인문/사회계열 110명(33.3%), 예체능계열 96명(29.1%), 이공/자연계열 86명(26.1%), 의약학계열 38명(11.5%)으로 분포하였다. 활용 빈도에 있어서는 주 1~2회 이상 사용하는 비율이 93.3%에 달하여 대학생들의 AI 활용이 보편화 되었음을 확인하였다.

표 1. 측정 변수의 조작적 정의 및 문항 출처

Table 1. Operational definitions of measurement variables and sources of items

Latent Variable	Op. Def. & Meas. Content	Items	Scale	Ref. Meas. Tool (Rev.)
AI Literacy	Understanding AI principles, recognizing hallucinations, prompt manipulation, and truth verification	5	5-pt	[17] [18] [19]
Result Dependence	Cognitive offloading, uncritical copying of AI answers, and total reliance	2	5-pt	[10] [8] [9]
Critical Interaction	Proactive collaboration as a critical partner for idea refinement and logical error detection	2	5-pt	[14] [12] [29]
Learning Well-being	Relief of task stress from loss of control, academic satisfaction, and positive psychological state	3	5-pt	[16] [32] [15]
Ethical Sensitivity	Vigilance against plagiarism, valuing evaluation fairness, and academic integrity	5	5-pt	[22] [24] [33]

표 2. 조사 대상자의 일반적 특성(N=330)

Table 2. General characteristics of survey respondents (N=330)

Category	Characteristics	Frequency (n)	Percentage (%)	Cumulative %
Gender	Male	168	50.9	50.9
	Female	162	49.1	100.0
Grade	1st Year	150	45.5	45.5
	2nd Year	97	29.4	74.9
	3rd Year	41	12.4	87.3
	4th Year or above	42	12.7	100.0
Major Field	Humanities/Social Sciences	110	33.3	33.3
	STEM	86	26.1	59.4
	Arts/Physical Education	96	29.1	88.5
	Medical/Pharmacy	38	11.5	100.0
AI Usage Frequency	Less than once a week	22	6.7	6.7
	1-2 times a week	144	43.6	50.3
	3-4 times a week	69	20.9	71.2
	5-6 times a week	49	14.8	86.0
	Daily (Constant access)	46	13.9	100.0
Total		330	100.0	

IV. 실증분석 결과

4-1 공통방법편의 및 신뢰도·집중 타당도 검증

자기보고식 설문 의 한계인 공통방법편의(CMB)를 확인하기 위해 Harman의 단일요인 검증을 실시한 결과는 표 3과 같다.

가장 큰 단일요인의 설명력이 34.2%로 나타나 기준치

(50% 미만)를 충족하였다. 측정모형의 확인적 요인분석(CFA) 결과, 적합도 지수는 $\chi^2=201.55(df=109)$, χ^2/df (수용합치도)=1.849, CFI(중분적합지수)=.955, TLI(중분적합지수)=.944, RMSEA(절대적합지수)=.048로 우수하게 도출되었다. 분석결과는 표 3과 같다. 모든 측정 문항의 표준화 요인적재량(Standardized λ)은 .70 이상으로 유의($p<.001$)하였으며, 복합신뢰도(CR)는 .820 이상, 평균분산추출지수(AVE)는 .615 이상으로 나타나 Hair 등이 제시한[35] 기준 요건을

표 3. 측정도구의 요인적재량, 신뢰도 및 집중 타당도 분석 결과

Table 3. Factor loading, reliability, and convergent validity analysis results of measurement tools

Latent Variable	Item	EFA Loading	CFA Std. λ	S.E.	C.R. (t)	Cronbach's α	CR	AVE
AI Literacy	Q1	.780	.795	0.045	17.666***	.884	.886	.615
	Q2	.821	.812	0.048	18.152***			
	Q3	.851	.860	0.051	19.340***			
	Q4	.812	.825	0.046	18.520***			
	Q5	.795	.804	0.044	17.890***			
Result Dependence	Q6	.815	.822	0.052	15.808***	.818	.820	.695
	Q7	.830	.845	0.055	16.210***			
Critical Interaction	Q8	.821	.835	0.049	16.890***	.825	.828	.708
	Q9	.842	.858	0.051	17.102***			
Ethical Sensitivity	Q10	.810	.825	0.048	17.880***	.905	.908	.665
	Q11	.852	.868	0.052	18.995***			
	Q12	.872	.885	0.055	19.445***			
	Q13	.765	.780	0.044	16.220***			
	Q14	.788	.802	0.046	16.850***			
Learning Well-being	Q15	.758	.762	0.048	15.220***	.835	.838	.634
	Q16	.792	.805	0.050	16.110***			
	Q17	.815	.825	0.052	16.550***			

충족하여 집중 타당성이 입증되었다.

4-2 변수 간 상관관계 및 판별 타당도 분석

Fornell과 Larcker의 기준[36]에 따라 판별 타당도를 검증한 결과는 표 4와 같다. 각 잠재변수의 AVE 제곱근 값(대각선 굵은 글씨)이 해당 변수와 타 변수 간의 상관관계수 최댓값(.632)보다 높게 나타나 제 변수 간의 판별 타당성이 확보되었다.

표 4. 잠재변수 간 상관관계 행렬 및 판별 타당도

Table 4. Correlation matrix and discriminant validity between latent variables

Variable Name	1	2	3	4	5
1. AI Literacy	.784				
2. Result Dependence	-.315***	.833			
3. Critical Interaction	.632***	-.128*	.841		
4. Learning Well-being	.485***	-.245***	.588***	.796	
5. Ethical Sensitivity	.352***	-.402***	.415***	.382***	.815

Note: Diagonal elements in bold represent the square root of the AVE for each latent variable; off-diagonal elements are correlations. (* $p < .05$, ** $p < .01$, *** $p < .001$)

표 5. 생성형 AI 활용 빈도에 따른 집단 간 차이분석 (ANOVA)

Table 5. Analysis of group differences according to generative AI usage frequency (ANOVA)

Variable Name	Usage Frequency (week)	N	Mean (M)	SD	F	p
AI Literacy	(1) Less than once	22	2.95	1.05	7.026	***
	(2) 1-2 times	144	3.42	0.81		
	(3) 3-4 times	69	3.83	0.82		
	(4) Near daily	95	3.75	0.88		
Critical Interaction	(1) Less than once	22	2.30	1.13	10.801	***
	(2) 1-2 times	144	3.30	1.05		
	(3) 3-4 times	69	3.70	0.93		
	(4) Near daily	95	3.73	1.02		
Learning Well-being	(1) Less than once	22	2.18	0.84	19.759	***
	(2) 1-2 times	144	3.62	0.90		
	(3) 3-4 times	69	3.93	0.88		
	(4) Near daily	95	4.06	0.91		
Ethical Sensitivity	(1) Less than once	22	3.66	1.03	0.585	.674
	(2) 1-2 times	144	3.69	0.85		(n.s.)
	(3) 3-4 times	69	3.83	0.84		
	(4) Near daily	95	3.73	0.90		

4-3 생성형 AI 활용 빈도에 따른 차이분석 (ANOVA)

구조모형 분석에 앞서 AI 활용 빈도가 주요 변수에 미치는 영향을 분석하기 위해 일원배치 분산분석(ANOVA)을 실시하였다. 분석 결과는 표 5와 같다.

AI를 빈번하게 사용할수록 'AI 리터러시($F=7.026$)', '비판적 상호작용($F=10.801$)', '학습 웰빙($F=19.759$)' 수준은 통계적으로 유의하게 상승하였다($p < .001$). 그러나 본 연구의 조절변수인 '윤리적 민감성'은 기술 활용 빈도와 관계없이 집단 간 유의미한 차이가 존재하지 않았다($F=0.585$, $p=.674$). 이는 집단 간 단순 평균 차이(ANOVA)는 나타나지 않았으나, 구조방정식 모델의 상호작용항 검증을 통해 특정 변수(윤리적 민감성)가 경로의 영향력을 변화시키는 조절 기제는 유의미하게 작동함을 확인하였다.

4-4 구조방정식 모형(SEM)을 통한 경로 분석 및 완전매개효과 검증

구조모형의 전반적 적합도는 $\chi^2/df=2.01$, CFI=.948, TLI=.938, RMSEA=.050으로 양호하게 나타나 가설 검증 기준을 충족하였다. 먼저 구조 모형의 각 인과 경로계수를 검증한 결과는 표 6과 같다.

AI 리터러시는 '결과 의존도'를 유의미하게 억제($\beta=-.320$, $p < .001$)하고 '비판적 상호작용'에 유의미한 정(+)의 영향을 미쳤다($\beta=.642$, $p < .001$)(가설 1a, 1b 채택). 또한, 인지적 오프로딩인 '결과 의존도'는 웰빙을 유의하게 저하($\beta=-.215$, $p < .001$)시킨 반면, '비판적 상호작용'은 웰빙을 유의미하게 향상($\beta=.465$, $p < .001$)시켰다(가설 2a, 2b 채택).

주목할 점은, 가설 3a에서 예측했던 바와 달리 AI 리터러시가 학습 웰빙에 미치는 직접효과(Direct Effect)는 통계적으로 유의하지 않은 것으로 나타나 기각되었다($\beta=.045$, $p=.321$).

나아가 매개효과의 유의성과 형태를 명확히 규명하기 위해 부트스트래핑(Bootstrapping, $N=5,000$)을 실시하여 총효과(Total Effect)를 분해하였다. 분석 결과는 표 7과 같다. AI 리터러시가 학습 웰빙에 미치는 총 효과($\beta=.412$)는 직접 효과($\beta=.045$)와 두 매개 경로를 통한 간접 효과의 총합($\beta=.367$)으로 구성됨을 확인하였다. AI 리터러시가 '결과 의존도'를 거쳐 웰빙에 미치는 간접효과($\beta=.069$)와 '비판적 상호작용'을 거쳐 웰빙에 미치는 간접효과($\beta=.298$)는 95% 신뢰구간에서 영(0)을 포함하지 않아 모두 통계적으로 유의미하였다(가설 3b, 3c 채택). 앞서 가설 3a에서 확인한 바와 같이 직접효과는 유의하지 않았으므로, 본 연구의 활용 방식(결과 의존도, 비판적 상호작용)은 AI 리터러시와 학습 웰빙 사이에서 '완전매개(Full Mediation)' 역할을 수행함이 최종적으로 입증되었다.

표 6. 구조방정식 모형 가설 검증 결과 (직접 인과 경로)

Table 6. SEM hypothesis testing results (direct causal paths)

Hypo.	Causal Path	B	β	S.E.	C.R. (t)	p	Result
H1a	AI Literacy → Result Dependence	-0.355	-0.320	0.056	-6.339	***	Accepted
H1b	AI Literacy → Critical Interaction	0.625	0.642	0.051	12.254	***	Accepted
H2a	Result Dependence → Learning Well-being	-0.192	-0.215	0.045	-4.266	***	Accepted
H2b	Critical Interaction → Learning Well-being	0.448	0.465	0.048	9.333	***	Accepted
H3a	AI Literacy → Learning Well-being (Direct)	0.042	0.045	0.042	1.000	.321	Rejected

표 7. 효과 분해(Effect Decomposition) 및 매개효과 부트스트래핑 검증 결과

Table 7. Effect decomposition and mediation effect bootstrapping test results

Path (AI Literacy → Well-being)	β	95% C.I. (Lower~Upper)	Result
Total Effect	.412**	(.315 ~ .520)	Accepted
Direct Effect	.045 (n.s.)	(-.042 ~ .132)	Rejected (Full Mediation)
Total Indirect Effect	.367***	(.265 ~ .468)	Accepted
H3b: via Result Dependence	.069*	(.012 ~ .135)	Accepted
H3c: via Critical Interaction	.298***	(.204 ~ .382)	Accepted

4-5 조절효과 분석: 윤리적 민감성의 상호작용 검증

조절효과 분석을 위해 투입된 독립변수와 조절변수는 다중 공선성 방지를 위해 모두 평균 중심화(Mean Centering)되었다. AMOS 내 자체 연산의 한계를 보완하기 위해 파이썬(Python 3.x) 환경에서 전체 표본 평균을 기반으로 편차 점수 및 상호작용항($X \times W$)을 생성한 후, 이를 모델에 투입하는 2단계 접근법을 취하였다. 본 분석은 관측변수의 평균점수를 활용한 경로분석 모델로 추정되었다.

단일 조절변수인 학습자의 윤리적 민감성이 유발하는 이중 조절효과를 검증한 결과는 표 8과 같다. 첫째, [결과 의존도 $\times$ 윤리적 민감성] 상호작용항은 학습 웰빙에 통계적으로 유

의미한 부(-)의 조절효과($\beta = -.292, p < .001$)를 보였다. 이는 학습 윤리가 투철한 학생일수록 생성형 AI에 의존했을 때 심리적 갈등이 가중되어 웰빙의 유의미한 저하가 더욱 심화됨을 통계적으로 지지하며(가설 4a 채택), 본 연구에서 가정된 ‘윤리적 부조화’ 현상의 발생을 입증한다.

둘째, [비판적 상호작용 $\times$ 윤리적 민감성] 상호작용항은 학습 웰빙에 유의미한 정(+)의 조절효과($\beta = .245, p < .001$)를 나타냈다. 도덕적 잣대가 높은 학생이 학업적 진실성을 잃지 않고 비판적 협업을 수행해 냈을 때 웰빙의 유의미한 증진이 가속화됨을 확인하였으며(가설 4b 채택), 이는 본 연구가 예상한 ‘윤리적 시너지’ 기제가 작용함을 시사한다.

조절효과의 구체적인 상호작용 양상을 시각적으로 확인하기 위해, 조절변수인 윤리적 민감성의 평균을 기준으로 집단을 구분(± 1 표준편차)하여 단순기울기 분석(Simple Slope Analysis)을 실시하고 이를 그림 2와 같이 도식화하였다.

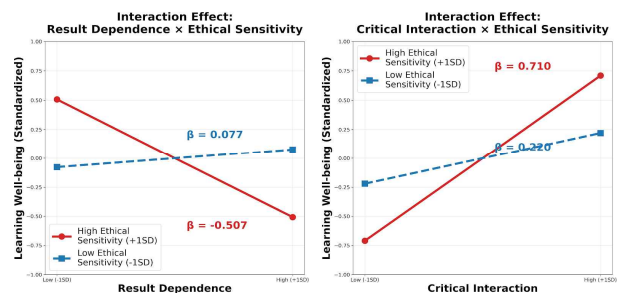


그림 2. 윤리적 민감성의 조절효과

Fig. 2. Moderating effects of ethical sensitivity

표 8. 윤리적 민감성의 상호작용(조절) 효과 분석 결과

Table 8. Analysis results of interaction (moderation) effects of ethical sensitivity

Hypo.	Interaction Path	B	β	S.E.	C.R. (t)	p	Result
Ctrl	Ethical Sensitivity → Well-being	0.220	0.235	0.044	5.000	***	Significant
H4a	Result Dep. $\times$ Ethical Sensitivity → Well-being	-0.218	-0.292	0.046	-4.739	***	Accepted
H4b	Critical Int. $\times$ Ethical Sensitivity → Well-being	0.175	0.245	0.040	4.375	***	Accepted

그림 2의 왼쪽 그래프에서 나타나듯, 결과 의존도가 학습 웰빙에 미치는 부(-)의 영향은 학습자의 윤리적 민감성에 의해 유의미하게 조절되었다. 윤리적 민감성이 낮은 집단(-1SD)에 비해 윤리적 민감성이 높은 집단(+1SD)에서 그 하락 기울기가 훨씬 더 가파르게 나타났다. 이는 윤리적 책임감이 강한 학생일수록 AI에 맹목적으로 의존할 때 학업적 불안과 죄책감이 가중되어 웰빙이 급격히 저하되는 ‘윤리적 부조화(Ethical Dissonance)’ 현상을 시각적으로 명확히 입증한다.

반면, 그림 2의 오른쪽 그래프를 살펴보면, 비판적 상호작용이 학습 웰빙에 미치는 정(+)의 영향 역시 윤리적 민감성에 의해 촉진되는 방향으로 조절됨을 알 수 있다. 윤리적 민감성이 낮은 집단(-1SD)보다 높은 집단(+1SD)에서 상승 기울기가 현저히 가파르게 나타났다. 즉, 높은 윤리 의식을 지닌 학생이 AI 시스템의 결과물을 맹신하지 않고 주도적으로 검증 및 확장해 낼 때, 학자적 양심을 지켜냈다는 내적 정당성이 더해져 학업적 웰빙이 극대화되는 ‘윤리적 시너지(Ethical Synergy)’ 현상이 발현됨을 강력하게 지지하는 결과이다.

V. 논의(Discussion)

본 연구는 대학생의 생성형 AI 리터러시가 활용 방식(결과 의존도, 비판적 상호작용)을 매개하여 학습 웰빙에 미치는 구조적 영향을 규명하고, 이 경로에서 학습자의 윤리적 민감성이 유발하는 이중 조절효과를 실증하였다. 도출된 연구 결과를 이론적 배경에서 고찰한 선행연구의 맥락과 연결하여 논의하면 다음과 같다.

5-1 완전매개 구조 확인과 인지적 오프로딩의 한계

본 연구의 가장 두드러진 발견 중 하나는 효과 분해를 통해 AI 리터러시가 학습 웰빙에 미치는 영향이 ‘완전매개(Full Mediation)’ 구조를 지닌다는 점을 실증한 것이다. 구조모형 분석 결과, AI 리터러시가 학습 웰빙에 정(+)의 직접효과를 미칠 것이라는 가설 3a는 기각되었다($\beta=.045$, $p>.05$). 오직 인지적 오프로딩인 ‘결과 의존도’를 억제하거나(가설 3b 채택), 주도적인 ‘비판적 상호작용’을 거칠 때(가설 3c 채택)에만 웰빙의 유의미한 증진으로 이어지는 강력한 간접효과가 확인되었다.

이러한 완전매개 구조는 학술적으로 매우 중대한 시사점을 제공한다. 학습자가 AI 시스템의 작동 원리나 할루시네이션의 위험성을 단순한 지식으로 이해(리터러시)하는 것 자체만으로는 학업적 스트레스나 불안을 해소할 수 없음을 뜻한다. 리터러시 역량이 실제 과제 수행 과정에서 생성형 AI에 대한 결과 의존[8],[30]을 거부하고 AI 시스템과 능동적으로 상호작용하는 ‘실천적 행동’으로 발현될 때 비로소 학습 웰빙이라

는 내적 보상을 획득할 수 있음을 증명한 것이다[13]. 이는 초기 인지적 과부하를 덜어줌과 동시에 문제 해결의 주도권을 학습자에게 귀속시키는 비판적 전유의 과정이 필수적임을 의미한다[7]. 즉, 생성형 AI 환경에서 단순한 기술적 지식보다, 기술을 어떻게 통제할 것인가에 대한 실질적인 ‘활용 궤적’이 심리적 성취를 결정짓는 핵심 기제임이 검증되었다.

5-2 규범적 이탈에서 심리적 갈등으로: ‘윤리적 부조화’ 현상의 규명

가설 4a를 통해 도출된 [결과 의존도 $\times$ 윤리적 민감성]의 유의미한 부(-)의 상호작용 효과는, 고등교육 현장의 AI 의존 현상을 단순히 학업적 부정행위(Academic Dishonesty)나 윤리적 이탈의 문제로만 환원해 온 기존 문헌[25],[20]의 접근에 인식의 전환을 요구한다. 분석 결과, 표절이나 무단 도용을 경계하는 도덕적 잣대가 높은 학생일수록 시간 부족이나 성적 압박 등으로 불가피하게 생성형 AI에 의존하게 되었을 때 학습 웰빙의 저하가 유의미하게 심화되는 현상이 확인되었다.

이는 AI 과의존을 징벌적 시각에서 벗어나, 양심적인 학습자가 겪는 심리적 갈등이라는 기제를 실증해 냈다는 데 의미가 있다. Festinger의 인지적 부조화 이론[23]에 비추어 볼 때, 이들은 학자적 양심을 수호하고자 하는 ‘내적 신념’과, 평가의 압박 속에서 사유의 통제권을 AI 시스템에 위임해 버린 ‘실제 행동’ 사이에서 심리적 죄책감을 경험하게 된다[24]. 턴잇인(Turnitin)과 같은 사후 적발 시스템에 대한 두려움[31]과 스스로 학업적 진실성을 훼손했다는 자괴감이 결합하여 발현되는 이 ‘윤리적 부조화(Ethical Dissonance)’ 현상은, 통제 지향적 정책이 자칫 학업 윤리가 투철한 학생들의 심리적 스트레스를 가중시키는 부작용을 낳을 수 있음을 시사한다.

5-3 도덕적 주체성의 회복과 ‘윤리적 시너지’ 현상의 발현

가설 4b를 통해 실증된 [비판적 상호작용 $\times$ 윤리적 민감성]의 정(+)의 조절효과는 생성형 AI 생태계에서 도덕적 주체성(Moral Agency)이 지니는 가치를 입증한다. 위험 회피에 머물던 수동적 윤리 연구[22]의 관점을 넘어, 도덕적 기준이 비판적 상호작용과 결합할 때 학습자에게 ‘학업적 정당성’을 부여하는 매커니즘으로 작용함이 확인되었다.

최현실과 윤영돈이 제안[33],[34]한 바와 같이, 도덕적 주체로서의 고유성을 잃지 않고 AI 산출물의 오류를 스스로 통제해 낸 학습자는 인간-AI 협력 관계를 주도적으로 이끌었다는 보상을 얻는다. 학업 윤리가 투철한 학생이 AI가 산출한 초안을 수동적으로 모방하지 않고 비판적 과정을 거쳐 지식 산출의 무결성(Integrity)을 지켜냈을 때, 이들은 혁신 기술의 편익과 학자적 양심 사이의 균형[37]을 달성하게 된다. 이렇

게 획득된 학업적 정당성은 단순한 과제 제출의 만족감을 뛰어넘어, 학습 웰빙의 유의미한 증진을 견인하는 ‘윤리적 시너지(Ethical Synergy)’ 기제로 기능함을 본 연구에서 실증적으로 뒷받침한다.

종합하자면, 본 연구는 기존 문헌들이 주로 다루었던 단편적인 AI 활용 성과나 단순한 기술적 수용 모델을 넘어, 학습자의 인지적·윤리적 기제를 복합적으로 규명하였다는 점에서 학술적 의의를 지닌다. 지금까지 논의된 바를 바탕으로, 관련 선행연구와 본 연구의 핵심적인 접근법 차이 및 학술적 확장성을 요약하면 표 9와 같다.

표 9. 선행연구 대비 본 연구의 주요 접근법 비교

Table 9. Comparison of key approaches between prior research and the current study

Research Dimension	Prior Research Trends	The Current Study's Perspective	Academic Significance & Ref.
Mediating Mechanism	Primarily focused on single-path models (e.g., AI usage frequency or self-efficacy).	Examines dual mediating paths: Result Dependence vs. Critical Interaction.	Clarifies the dual cognitive nature (offloading vs. engagement) of AI adoption[30],[31]
Moderating Factors	Mostly focused on demographic variables (e.g., gender, major, or grade).	Introduces Ethical Sensitivity as a psychological and value-based moderator.	Demonstrates how internalized values moderate the psychological outcomes of technology use[24],[30]
Outcome Variable	Measured through academic performance or general user satisfaction.	Focuses on Learning Well-being (Flourishing) as the ultimate goal.	Shifts the paradigm towards long-term learner prosperity in the AI era[16].

VI. 결론 및 제언

6-1 연구 결과 요약

본 연구는 330명의 대학생 표본을 바탕으로 조절된 이중 매개모형(Dual Moderated Mediation Model)을 구축하여, 생성형 AI 리터러시가 상반된 활용 방식을 거쳐 학습 웰빙에 미치는 영향과 단일 조절변수인 윤리적 민감성의 상호작용 효과를 실증하였다.

분석 결과 첫째, AI 리터러시는 학습 웰빙에 유의미한 정(+)의 직접효과를 나타내지 않았으나, 인지적 오프로딩인 ‘결과 의존도’를 억제하고 ‘비판적 상호작용’을 유의미하게 촉진하는 선행 변인으로 작용하였다. 둘째, AI 결과 의존도는 학습 웰빙을 훼손한 반면 비판적 상호작용은 웰빙을 향상시켰으며, 이를 통해 AI 리터러시가 행동적 상호작용을 거쳐야만

웰빙으로 이어지는 ‘완전매개’ 구조가 입증되었다. 셋째, 조절 효과 분석 결과, 학습 윤리가 투철한 학생이 생성형 AI에 맹목적으로 의존할 때 심리적 죄책감이 가중되어 웰빙 저하가 심화되는 ‘윤리적 부조화’ 현상이 유의미하게 나타났다. 반면, 이들이 주도적인 비판적 상호작용을 수행할 때는 학업적 정당성을 부여받아 웰빙 증진이 가속화되는 ‘윤리적 시너지’ 현상이 검증되었다.

6-2 실무적 및 정책적 시사점

본 연구의 결과는 대학교육 현장에서 생성형 AI를 단순한 도구가 아닌 학습자의 성장을 돕는 파트너로 정착시키기 위한 실천적 가이드라인을 제공한다. 구체적인 대학교육 적용 방안은 다음과 같다.

첫째, 대학은 기존의 기능 중심 AI 교육에서 벗어나 ‘비판적 검증 역량’을 강화하는 정규 커리큘럼을 구축해야 한다. 단순히 프롬프트 작성법을 가르치는 것에 그치지 않고, AI 산출물의 할루시네이션(Hallucination)과 편향성을 학생들이 직접 추적하고 교정하는 ‘AI 팩트체크 실습’을 전공 및 교양 필수 과정에 도입할 필요가 있다. 이는 학습자가 AI에 맹목적으로 의존하지 않고 주도적인 ‘비판적 상호작용’을 수행하게 함으로써 학업적 자기효능감과 웰빙을 증진시키는 토대가 된다.

둘째, 학습자의 윤리적 민감성을 실천적으로 함양하기 위한 ‘AI 학술 무결성(Academic Integrity) 교육 프로그램’의 의무화가 요구된다. 대학 차원에서 AI 활용 범위를 명확히 규정하는 가이드라인을 제시하고, 학생들이 자신의 과제물에서 AI의 기여도를 투명하게 밝히는 ‘AI 출처 표기법’을 체득하게 해야 한다. 이러한 교육은 본 연구에서 확인된 ‘윤리적 부조화’를 방지하고, 학생들이 도덕적 주체로서 기술을 통제하고 있다는 ‘윤리적 시너지’를 경험하게 하는 핵심 기제가 될 것이다.

셋째, ‘AI 윤리 샌드박스(Ethics Sandbox)’ 형태의 체험형 교육 환경 조성이 필요하다. 학생들이 안전한 환경에서 AI의 윤리적 딜레마 사례를 시뮬레이션해 보고, 기술 오남용이 학습 안녕감에 미치는 부정적 영향을 성찰해 보는 토론식 수업을 강화해야 한다. 이러한 고등교육 차원의 다각적인 개입은 대학생들이 급변하는 기술 환경 속에서 매몰되지 않고, 주체적이고 윤리적인 태도로 자신의 학습 웰빙을 지켜나가는 역량을 확보하는 데 기여할 것이다

6-3 연구의 한계 및 향후 연구 방향

본 연구는 학술적 기여에도 불구하고 다음과 같은 한계를 지닌다.

첫째, 횡단적(Cross-sectional) 설문 데이터를 통해 단일 시점에서의 인과관계를 측정하였으므로, 향후 종단적(Longitudinal) 설계를 적용하여 AI 수용 기간에 따른 학습 웰빙의 장기적 변화 궤적을 추적할 필요가 있다.

둘째, 자가보고식(Self-report) 설문 특성상 응답자의 실제 학업 성취도(학점 등)를 반영하는 데 한계가 있다. 따라서 후속 연구에서는 학생들의 실제 성적 데이터나 튜터링 프로그램의 성과 지표[5]를 결합한 객관적 성과(Performance) 분석이 수반된다면 연구의 타당성을 더욱 공고히 할 수 있을 것이다.

셋째, 본 연구의 주요 매개변수인 ‘결과 의존도’와 ‘비판적 상호작용’이 한정된 설문 환경을 고려하여 2문항 척도로 측정되었다는 점이다. 비록 선행연구의 핵심 지표를 추출하여 내용 타당도와 통계적 신뢰도를 확보하였으나, 향후 연구에서는 해당 구성 개념을 보다 다차원적으로 포착할 수 있는 다문항 척도(Multi-item scale)를 개발하고 교차 타당성을 검증하는 작업이 요구된다.

참고문헌

- [1] S. Zafar, F. Shaheen, and J. Rehan, “Use of ChatGPT and Generative AI in Higher Education: Opportunities, Obstacles and Impact on Student Performance,” *IRASD Journal of Educational Research*, Vol. 5, No. 1, pp. 1-12, December 2024. <https://doi.org/10.52131/jer.2024.v5i1.2463>
- [2] X. Zhai, X. Chu, C. S. Chai, M. S. Y. Jong, A. Istenic, M. Spector, ... and Y. Li, “A Review of Artificial Intelligence (AI) in Education from 2010 to 2020,” *Complexity*, Vol. 2021, pp. 1-21, 2021. <https://doi.org/10.1155/2021/8812542>
- [3] H. Lee and J. W. Yoo, “Exploration of University Students’ Educational Use Experiences and Perceptions of Generative AI: A Case Study of A University,” *Journal of the Korea Contents Association*, Vol. 24, No. 1, pp. 428-437, January 2024. <https://doi.org/10.5392/JKCA.2024.24.01.428>
- [4] J.-Y. Yeom and I. Park, “Exploring the Acceptance Process of ChatGPT Among University Students Based on the Theory of Planned Behavior,” *Journal of Korean Association for Educational Information and Media*, Vol. 30, No. 1, pp. 1-26, 2024. <http://doi.org/10.15833/KAFEIA M.30.1.001>
- [5] M. Kim and Y.-H. Jang, “Effects of Tutoring Programs on Core Competencies of University Students and Capacity Difference Analysis by Student Variables,” *Journal of Lifelong Learning Society*, Vol. 17, No. 3, pp. 267-290, 2021. <https://doi.org/10.26857/JLLS.2021.8.17.3.267>
- [6] J. Y. Kim, Impact of Generative AI Service Characteristics, Perceived Utility, and Sacrifice on Perceived Value and Satisfaction: A Study Applying the Value-Based Acceptance Model for ChatGPT Users, Ph.D. Dissertation, Hansung University, Seoul, 2025.
- [7] M. Gerlich, “From Offloading to Engagement: An Experimental Study on Structured Prompting and Critical Reasoning with Generative AI,” *Data*, Vol. 10, No. 11, 172, 2025. <https://doi.org/10.3390/data10110172>
- [8] S. Zhang, X. Zhao, T. Zhou, and J. H. Kim, “Do you Have AI Dependency? The Roles of Academic Self-Efficacy, Academic Stress, and Performance Expectations on Problematic AI Usage Behavior,” *International Journal of Educational Technology in Higher Education*, Vol. 21, 34, 2024. <https://doi.org/10.1186/s41239-024-00467-0>
- [9] Z. Revesai, “Generative AI Dependency: The Emerging Academic Crisis and Its Impact on Student Performance—A Case Study of a University in Zimbabwe,” *Cogent Education*, Vol. 12, No. 1, 2549787, 2025. <https://doi.org/10.1080/2331186X.2025.2549787>
- [10] J.-S. Seo and S.-Y. Lim, “The Influence of ChatGPT Use on University Students’ Critical Thinking Skills,” *Journal of Families and Better Life*, Vol. 43, No. 3, pp. 109-120, 2025. <https://doi.org/10.7466/JFBL.2025.43.3.109>
- [11] Y. M. Kim, “The Effect of Performing Projects Using ChatGPT on the Higher-Order Thinking Skills of University Freshmen,” *Journal of General Education*, Vol. 39, pp. 121-151, 2024. <http://doi.org/10.31658/DSHR.39.5>
- [12] J. Yoo and S. Kim, “Analysis of Students’ Interaction Experiences in Problem Solving through Human-Generative AI Collaboration,” *Journal of Practical Engineering Education*, Vol. 18, No. 1, pp. 67-76, 2026.
- [13] K. Yoo and S. Ahn, “The Impact of Generative AI-Integrated Education on Academic Self-Efficacy, Metacognition, and Problem-Solving Skills of University Students,” *Journal of Korean Association of Computer Education*, Vol. 27, No. 4, pp. 13-20, 2024. <https://doi.org/10.32431/kace.2024.27.4.002>
- [14] O. Y. Han, “Generative AI and Critical Thinking in Engineering Education: A Study on the Use of Prompt Engineering,” *Journal of Engineering Education Research*, Vol. 27, No. 6, pp. 38-47, 2024. <https://doi.org/10.18108/jeer.2024.27.6.38>
- [15] J. Liang, L. Wang, J. Luo, Y. Yan, and C. Fan, “The Relationship Between Student Interaction with Generative Artificial Intelligence and Learning Achievement: Serial Mediating Roles of Self-Efficacy and Cognitive Engagement,” *Frontiers in Psychology*, Vol. 14, 1285392, 2023. <http://doi.org/10.3389/fpsyg.2023.1285392>
- [16] H.-R. Cho, S.-J. Lee, S.-Y. Lee, and Y.-J. Kim, “Factors Affecting Satisfaction with the Application of Generative AI in Higher Education,” *Journal of the Korea Society of Computer and Information*, Vol. 29, No. 9, pp. 299-306,

2024. <https://doi.org/10.9708/jksoci.2024.29.09.299>
- [17] G. Y. Kang, Development of AI Literacy Measurement Tools for Undergraduate and Graduate Students, Ph.D. Dissertation, Yonsei University, Seoul, 2024.
- [18] E. K. Oh, "Analysis of University Students' AI Usage Characteristics and Perceptions," *Journal of the Korean Society for Library and Information Science*, Vol. 59, No. 1, pp. 671-692, 2024. <http://doi.org/10.4275/KSLIS.2025.59.1.671>
- [19] S. Y. Byeon, "Direction of AI Literacy Education in the Era of Generative AI," *Intelligence Information Ethics Issue Report*, Vol. 4, No. 4, pp. 1-13, 2023.
- [20] K. Bittle and O. El-Gayar, "Generative AI and Academic Integrity in Higher Education: A Systematic Review and Research Agenda," *Information*, Vol. 16, No. 4, 296, 2025. <https://doi.org/10.3390/info16040296>
- [21] A. K. Kofinas, C. H.-H. Tsay, and D. Pike, "The Impact of Generative AI on Academic Integrity of Authentic Assessments within a Higher Education Context," *British Journal of Educational Technology*, Vol. 56, No. 6, pp. 2522-2549, 2025. <https://doi.org/10.1111/bjet.13585>
- [22] N. Rizun, O. N. Bordean, A. Nikiforova, I. N. Beleiu, and A. Revina, "Generative AI in Higher Education: Ethical and Behavioral Factors Influencing Students' Intentions to Use ChatGPT," *Computers and Education Open*, Vol. 10, 100336, June 2026. <https://doi.org/10.1016/j.caeo.2026.100336>
- [23] L. Festinger, *A Theory of Cognitive Dissonance*. Stanford, CA: Stanford University Press, 1957.
- [24] S. Kang, "An Exploratory Analysis of University Students' Use of Generative AI Tools and Their Perceptions of Source Attribution and Ethical Compliance," *Journal of Educational Research*, Vol. 47, No. 2, pp. 103-124, 2025. <http://doi.org/10.35510/JER.2025.47.2.103>
- [25] E. Faisal, "Complementary AI in Higher Education: Behavioral, Cognitive, and Ethical Implications of ChatGPT and DeepSeek," *Frontiers in Psychology*, Vol. 17, 1699114, March 2026. <https://doi.org/10.3389/fpsyg.2026.1699114>
- [26] Y. Sun, D. Sheng, Z. Zhou, and Y. Wu, "AI Hallucination: Towards a Comprehensive Classification of Distorted Information in Artificial Intelligence-Generated Content," *Humanities and Social Sciences Communications*, Vol. 11, 1278, 2024. <https://doi.org/10.1057/s41599-024-03811-x>
- [27] H. S. Lee, "The Legal Responsibility Framework of Generative AI Hallucinations and Normative Design for the Protection of Fundamental Rights," *Legal Theory and Practice Review*, Vol. 14, No. 1, pp. 49-88, February 2026. <https://doi.org/10.30833/LTPR.2026.02.14.1.49>
- [28] Y. Kim and H.-S. Lee, "Effects of Generative AI Characteristics on User Satisfaction and Trust: AI Literacy as a Moderator," *Journal of the Korea Society of Computer and Information*, Vol. 30, No. 12, pp. 253-263, 2025. <https://doi.org/10.9708/jksoci.2025.30.12.253>
- [29] J. H. Heo, "A Study on the Poetic Appropriation of Mechanical Errors: Focusing on AI Hallucination and Poetic Language," *Journal of Korean Literary Criticism*, No. 87, pp. 39-69, 2025. <https://doi.org/10.35832/kmlc..87.202509.39>
- [30] S. Y. Erlangga, Sarwanto, and Harlita, "Unpacking the Cognitive and Ethical Pathways of Generative AI Tools in Higher Education: A PLS-SEM Study of Learning Performance Mediation and Moderation Effects," *Research in Learning Technology*, Vol. 34, 2026. <https://doi.org/10.25304/rlt.v34.3563>
- [31] G. Pitts, N. Rani, W. Mildort, and E.-M. Cook, "Students' Reliance on AI in Higher Education: Identifying Contributing Factors," arXiv:2506.13845, 2025. <https://doi.org/10.48550/arXiv.2506.13845>
- [32] W. Jia, L. Pan, and S. Neary, "Effect of GenAI Dependency on University Students' Academic Achievement: The Mediating Role of Self-Efficacy and Moderating Role of Perceived Teacher Caring," *Behavioral Sciences*, Vol. 15, No. 10, 1348, 2025. <https://doi.org/10.3390/bs15101348>
- [33] H. S. Choi, "Systematizing Ethics in the AI Era: 3-Layer Ethical Structure of Human-Machine Conditional Cooperation Through Trolley Dilemma Recognition Analysis," *Journal of Asian Studies*, Vol. 28, No. 4, pp. 437-454, 2025.
- [34] Y. D. Yoon, "2022 Revised Moral Curriculum and AI Ethics Education Direction - Focusing on the Uniqueness and Individuality of Moral Subjects," *Journal of Moral & Ethics Education*, No. 89, pp. 361-390, 2025. <https://doi.org/10.18338/kojmee.2025.89.361>
- [35] J. F. Hair, W. C. Black, B. J. Babin, and R. E. Anderson, *Multivariate Data Analysis*, 8th ed. London, UK: Cengage, 2019.
- [36] C. Fornell and D. F. Larcker, "Evaluating Structural Equation Models with Unobservable Variables and Measurement Error," *Journal of Marketing Research*, Vol. 18, No. 1, pp. 39-50, 1981. <https://doi.org/10.2307/3151312>
- [37] N. J. Francis, S. Jones, and D. P. Smith, "Generative AI in Higher Education: Balancing Innovation and Integrity," *British Journal of Biomedical Science*, Vol. 81, 2025. <https://doi.org/10.3389/bjbs.2024.14048>



김명한(Myoung-Han Kim)

1988년 : 조선대학교 경상대학

(경영학사)

2004년 : 연세대학교 경영대학원

(경영학석사)

2023년 : 호서대학교 벤처대학원

(경영학박사)

1990년~2021년: 국민은행, 서민주택금융재단

2022년~2025년: 조선대학교 미래사회융합대학 겸임교수

2023년~2025년: 유한대학교 기초융합교육원 강사

2022년~현 재: 김포대학교 경영학과 겸임교수

2025년~현 재: 동국대학교 행정대학원 대우교수

※관심분야 : 벤처경영, 디지털마케팅, 부동산자산관리, 매니지먼트