

컴퓨터교육학회 논문지 2026년 제29권 제6호
<https://doi.org/10.32431/kace.2026.29.6.011>



생성형 AI 활용의 학문적 진실성 판단을 위한 비네트 기반 도구 개발

Development of a Vignette-Based Instrument for Assessing Academic Integrity in the Use of Generative AI

박하나[†] 

Hana Park[†]

요약

본 연구의 목적은 생성형 AI 활용 맥락에서 학습자의 학문적 진실성 판단에 영향을 미치는 요인을 검토하고 교육적으로 활용할 수 있는 비네트 기반 도구를 개발하는 데 있다. 이를 위해 문헌연구를 통해 AI 활용의 학문적 진실성 판단 요인을 도출하고, 이를 기반으로 AI 활용의 학문적 진실성 판별과 교육을 위한 비네트를 개발하였다. 개발한 비네트는 전문가 검토와 사용자 검토를 통해 타당화하였다. 본 연구에서 개발된 비네트 기반 도구는 학습자의 학문적 진실성 판단을 진단하는 측정 도구이자, 상황 기반 토의와 성찰을 유도하는 교육적 도구로 활용될 수 있으며, AI 시대의 학문적 진실성 교육을 위한 실천적 기반을 제공한다는 점에서 의의를 가진다.

주제어 생성형 AI, AI 윤리, 학문적 진실성, 비네트

ABSTRACT

This study aims to examine the factors influencing learners' academic ethical judgments in the context of generative AI use and to develop a vignette-based instrument for educational application. To this end, a literature review was conducted to identify key factors related to academic integrity in AI use, and based on these factors, vignettes were developed to assess and support judgments of academic integrity in AI-assisted work. The developed vignettes were validated through expert review and user testing. The vignette-based instrument developed in this study can be used both as an assessment tool to diagnose learners' academic ethical judgments and as an instructional tool to facilitate scenario-based discussion and reflective thinking. This study is significant in that it provides a practical foundation for academic integrity education in the era of AI.

Keywords Generative AI, AI ethics, Academic Integrity, Vignette

†정회원 경기대학교 연구교수
 논문투고 2026년 04월 16일
 심사완료 2026년 05월 13일
 게재확정 2026년 05월 28일
 발행일자 2026년 06월 30일

1. 서론

학문적 진실성(academic integrity)이란 지식 생산과 공유의 전 과정에서 정직, 신뢰, 공정, 존중, 책임의 다섯 가지 핵심 가치를 준수하는 것을 의미한다[1]. 이는 연구결과의 신뢰성과 재현 가능성, 지적 소유권의 존중, 그리고 공정한 교육 환경을 유지하는 데 결정적인 역할을 하며, 지식 공동체의 존속과 발전을 가능하게 하는 구조적 기반이다. 교육기관과 학술단체는 학문적 진실성의 핵심 가치를 기초로 하여 공동체의 규범을 제도화 하였으며, 이러한 제도적 기준을 바탕으로 표절, 부정행위, 허위연구, 데이터 조작 등의 행위를 학문윤리의 위반으로 규정해왔다. 대표적인 학문적 부정행위인 표절의 경우, 타인의 저작물을 출처 없이 사용했는가의 여부가 판단의 핵심 기준이었다. 그러나 생성형 AI가 등장하면서 AI가 생성한 결과를 누구의 저작물로 보아야 하는가, 사용자는 실질적으로 어떤 기여를 했는가, AI 사용 사실을 고지해야 하는가와 같은 새로운 질문들이 등장하였다[2-6].

생성형 AI는 자료를 그대로 출력하는 것이 아니라, 재구조화하는 방식으로 산출물을 생성해낸다. 따라서 생성형 AI의 결과물은 기술적으로 원본에 해당하며, 기존 유사도 기반 표절 탐지 도구로는 판별에 한계가 따른다[7]. 또한 AI가 어떤 데이터를 학습하고, 어떤 자료를 기반으로 결과를 생성하였는지 파악이 어렵다는 특성으로 인해 사용자의 의도와 관계없이 표절 혹은 저작권 침해의 문제가 발생할 가능성이 있다.

생성형 AI 활용에 관한 학문적 부정행위의 경계는 단일한 기준으로 판단하기 어렵다. AI 활용 수준은 연속적 스펙트럼으로 분류하는 것이 일반적이며, 동일한 활용 방식도 맥락에 따라 다르게 평가될 수 있다. 대학생들의 인식을 조사한 결과에 따르면 AI의 생성물을 그대로 제출하는 것은 명확히 부정행위라고 인식한 반면, AI를 보조적으로 활용하거나, 아이디어 도출 등의 목적으로 활용한 경우에는 상황에 따라 다른 판단을 내리는 것으로 나타났다[8-10]. 그러나 여기에서도 양적, 질적으로 어디까지를 보조적 수단으로 볼 것인가에 대한 명확한 정의가 어렵다는 문제가 남는다.

이러한 판단의 복잡성은 AI 활용에 관여하는 요인이 복합적이라는 데서 기인한다. 선행연구에서 공통적으로 언급되는 요인으로는 AI와 사용자의 기여도, 투명성, AI 활용 맥락, 저작물에 대한 권리와 책임, 윤리성 등이 있다[2-6].

AI와 사용자 간 기여도에 대해 논할 때, 이는 단순히 AI가 생성한 비율이 얼마인가를 의미하지 않는다. 아이디어, 판단, 편집 등의 과정에서 학습자가 얼마나 실질적으로 개입했는지가 보다 중요하다[5, 11].

투명성은 전통적으로 학문윤리에서 중요하게 다루어진 요인 중 하나이다. 다수의 선행연구에서 AI 사용 미표기 자체를 기만적 행위로 규정하며, 사용 도구, 활용 방식, 기여 범위를 구체적으로 기술할 것을 권고한다[12, 13]. 그러나

여전히 많은 경우 명시적 기준이 부재하며, 이로 인하여 기관이나 교수자의 명시적 정책이 없는 상황에서도 암묵적 기대를 위반한 것으로 볼 수 있는지에 대한 논쟁이 지속되고 있다[14].

AI가 활용되는 맥락적 요인도 부정행위 판단에 중요하게 작용한다. 동일한 방식으로 AI를 사용하더라도, 과제의 목적에 따라 학습효율을 높이는 유용한 도구로 평가받기도 하고, AI에 의한 대리 수행으로 간주되어 학습목표를 저해한다는 비판을 받기도 한다[15].

이와 같은 맥락적 요인은 결과물에 대한 권리와 책임의 판단에도 영향을 미친다. 개인 사용자의 경우라면 프롬프트 엔지니어링의 창의적 기여를 어디까지 인정할 것인가가 쟁점이 될 수 있다[16]. 그러나 팀 프로젝트에서 일부 구성원이 AI를 사용한 경우, 이 쟁점은 책임을 팀원 전체에게 귀속시킬 수 있는가까지 확장된다. 이처럼 각각의 요인은 독립적으로 작용하기보다 상호작용하며 판단의 복잡성을 형성한다. AI의 생성 비율이 높더라도 투명성 요인을 충족하는 경우에는 판단이 달라질 수 있으며, 맥락에 따라 동일한 행위에 대한 판단 기준이 상이하게 적용될 수 있다. 이러한 복잡성으로 인하여 생성형 AI의 윤리적 활용을 위한 기준과 지침을 마련하는 것은 쉽지 않다.

이 같은 상황에서는 개인의 노력이 더욱 중요해진다. 다만 AI의 윤리적 활용을 위한 규범을 알고 있더라도 실제 상황에서는 앞서 살펴본 바와 같이 여러 요인들이 복합적으로 작용하여 경계가 모호해지기 때문에 판단이 쉽지 않다. 따라서 생성형 AI의 윤리적 활용을 위해서는 AI 부정행위의 개념을 이해하는 것이 아니라, 실제 맥락 속에서 부정행위 판별을 둘러싼 다양한 이슈와 논점을 직접 마주하고 성찰하는 경험이 필요하다. 동일한 상황도 어떤 요인을 중심으로 보느냐에 따라 판단이 달라질 수 있음을 체험하고, 자신의 판단 근거를 언어화하며, 타인의 관점과 비교하는 과정을 통해 윤리적 판단 역량이 형성될 수 있다[17].

특정 상황에 대한 판단이나 태도를 표명하도록 유도하기 위한 방법으로 비네트를 활용할 수 있다[18]. 실제 발생 가능한 상황을 간결하게 서술한 짧은 이야기를 의미하는 비네트는 추상적 윤리 규범을 구체적 맥락에서 탐색하는 데 적합하여 도덕적 판단 혹은 윤리적 딜레마를 다루는 연구와 교육에서 널리 활용되어 왔다[19, 20]. 특히 정답이 명확하지 않은 경계선 사례를 제시함으로써 응답자가 단순한 규범 확인이 아닌 비판적 추론 과정을 경험하도록 유도할 수 있다는 점에서 교육적 가치가 높다[21, 22]. 생성형 AI 활용의 다양한 경계선 사례는 단일한 판단 기준이 적용되기 어려운 복합적 현상이라는 점에서 비네트는 각 요인이 어떻게 작동하는지를 맥락 속에서 직접 탐색하고 자신의 판단을 점검하는 성찰적 교육 도구로서의 가능성을 갖는다. 이에 본 연구에서는 생성형 AI 활용의 학문적 진실성에 대한 윤리적 숙고와 성찰을 촉진하기 위한 교육도구로서 비네트 기반 질문지를 개발하고자 한다.

2. 연구방법

2.1 연구절차

본 연구는 Figure 1과 같은 절차로 진행되었다.

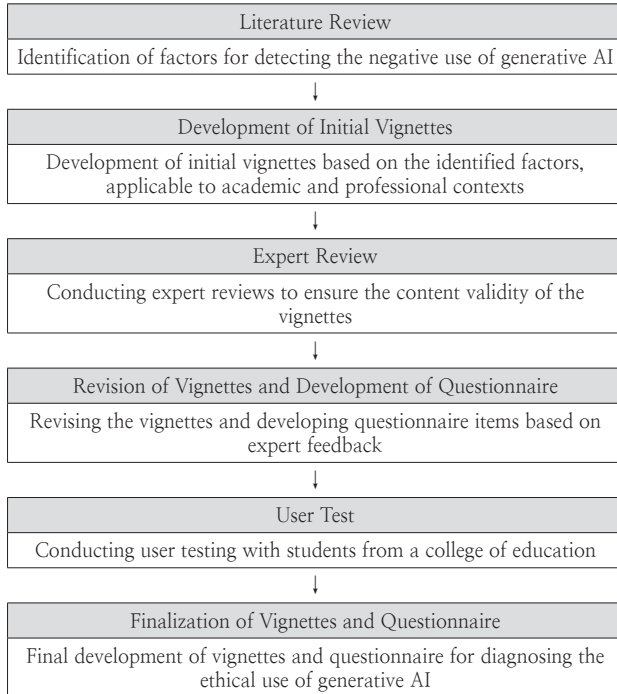


Figure 1. Research Procedure

먼저 생성형 AI 활용에 따른 학문적 부정행위의 판별 요인을 도출하기 위하여 다음과 같이 문헌분석을 실시하였다. 첫째, Web of Science와 Google scholar에서 ‘AI plagiarism’, ‘generative AI’와 ‘academic integrity’, ‘ChatGPT’와 ‘plagiarism’ 등의 키워드를 조합하여 최근 5년 간의 연구를 수집하였다. 둘째, 1차로 수집된 220편의 학술자료 중 중복자료를 포함하여 AI 탐지 기술의 기술적 성능만을 다룬 연구, 생성형 AI와 무관한 일반적 표절 탐지 연구를 제외하고, 77편을 최종 분석대상으로 선정하였다. 셋째, 최종 선정된 77편의 문헌을 대상으로 AI 기여도, 투명성, 맥락성, 저작성, 윤리성의 5가지 판별 요인을 도출하였다.

다음으로 문헌분석을 통해 도출한 판별요인을 반영하여 12개의 비네트를 개발하였다. 이를 위해 Table 1과 같이 과제 유형과 윤리적 쟁점을 포함한 비네트 설계안을 작성하고,

Claude를 통해 12개의 비네트 초안을 개발하였다. AI의 부정적 활용 판별 요인이 정확하고 충분하게 반영되었는지 검토하여 수정하였으며, 수정된 비네트는 전문가 검토를 통해 내용 타당도를 확보하였다.

전문가 검토 결과에 따라 수정한 2차 비네트는 교육대학원에 재학 중인 예비교사를 대상으로 비네트의 적절성 및 현실성을 평가하여 실제 학습 상황에서의 적용 가능성을 확인하였다.

2.2 비네트(Vignette) 개발

2.2.1 비네트 설계 원칙

비네트는 실제 발생 가능한 상황을 간결하게 서술한 짧은 이야기로, 응답자가 특정 상황에 대한 판단이나 태도를 표명하도록 유도하는 측정 도구이다[18, 20]. 도덕적 판단이나 윤리적 딜레마와 같이 추상적이고 민감한 주제를 구체적 맥락에서 탐색하는 데 적합하여 학문적 진실성 연구에서도 널리 활용되어 왔다[19, 22]. 본 연구에서는 도출된 5개 판별 요인을 비네트에 체계적으로 반영하기 위하여 다음과 같은 네 가지 설계 원칙을 수립하였다.

첫째, 판별 요인 대표성이다. 각 비네트는 5개 판별 요인 중 하나 이상을 핵심 요인으로 포함하며, 비네트 세트 전체 수준에서 요인 분포의 균형을 유지하도록 사전에 요인 매핑을 계획하였다.

둘째, 실제 맥락에서 부정행위의 판별은 복합적인 성격을 가진다는 점을 고려하여 핵심 요인과 주변 요인을 복합적으로 반영함으로써 판단 모호성을 확보하였다. 모호성이 낮은 사례부터 판단의 경계에 위치한 사례를 고루 포함함으로써 응답자가 AI 활용에 따른 부정행위의 판별이 갖는 복합적인 성격을 이해하고 비판적 추론 경험을 할 수 있도록 구성하였다. 또한 유사한 과제 맥락 내에서 생성형 AI의 개입 수준과 기능적 역할을 단계적으로 변주함으로써, 특정 판단 요소의 변화가 응답자의 윤리적 판단에 미치는 영향을 비교해가며 탐색할 수 있도록 설계하였다.

셋째, 상황 다양성이다. 생성형 AI 활용 맥락이 학업 상황과 업무 상황에 따라 판단 기준이 달라질 수 있다는 점을 반영하여, 과제, 학술 연구, 직무 보고서 등 다양한 상황을 포함하였다[23].

넷째, 현실성이다. 실제 발생 가능한 맥락을 결합하여 현실성을 높임으로써 응답자의 몰입을 유도하였으며, 실제 맥락으로의 전이를 촉진하고자 하였다.

Table 1. Design of Vignette

Case	Task context	AI Usage			Ethical Issue
		AI Utilization Scope	Learner Contribution	Disclosure Status	
1	Academic (Report)	Full delegation (unmodified submission)	None	Non-disclosure	Direct submission of AI-generated content without any modification: baseline for complete plagiarism
2	Academic (Report)	Full delegation + surface editing	Minimal (synonym substitution, paragraph rearrangement)	Non-disclosure	Ethical effect of synonym replacement and paragraph rearrangement on AI-generated content

Case	Task context	AI Usage			Ethical Issue
		AI Utilization Scope	Learner Contribution	Disclosure Status	
3	Academic (Essay)	Partial (~60% AI-generated)	~40% directly authored	Non-disclosure	Mixed contribution: 60% AI-generated + 40% learner's own perspective
4	Academic (Literature Review)	Partial (~50%) + partial disclosure	~50% directly authored	Partial disclosure (one-line footnote)	Adequacy of disclosure method and its mitigating effect on ethical judgment
5	Academic (Coding Assignment)	Full delegation (auto-generated code)	None	Non-disclosure	AI-generated code submitted as own work: explicit violation of 'written by oneself' requirement
6	Professional (Market Analysis)	Partial (~70% AI-generated)	~30% directly authored	Non-disclosure	AI use in professional context: unverified facts and non-disclosure to stakeholders
7	Research (English Abstract)	Language assistance (translation/editing only)	Full content authored directly by researcher	Non-disclosure	Ethical permissibility boundary of AI as a language tool in academic writing
8	Academic (Literature Review)	AI summarization (as reading substitute)	Structural organization directly authored	Non-disclosure	Writing a report based on AI-generated summaries without directly reading the original papers
9	Academic (Team Project)	Partial (one team member's section fully AI-generated)	Other team members' sections directly authored	Non-disclosure within team	Collective responsibility attribution and intra-team deception
10	Academic (Proposal)	Idea generation (AI-sourced concepts)	Elaboration and detailing directly authored	Non-disclosure	AI as source of creative ideas in a task evaluating originality and creativity
11	Academic (Instrument Development)	Full item generation (all items AI-generated)	Only 3 items revised	False representation ('researcher-developed')	False labeling of AI-generated items as 'researcher-developed': research instrument fabrication
12	Professional (Policy Proposal)	Full delegation + minimal editing	Minimal	Non-disclosure	AI substituting expert knowledge and professional judgment in policy recommendation

2.2.2 1차 비네트 도출

각 비네트는 맥락, AI 사용 방식, 핵심 판별 요인을 고려하여 구성하였다. 맥락 요인으로는 사례가 발생하는 상황과 구체적인 과업의 유형을 반영하였다. AI 사용 방식에서는 AI와 사용자의 기여도와 AI 사용 고지 여부를 포함하였다. 마지막으로 핵심 판별 요인으로는 선행연구를 통해 도출한 학문적 진실성 판별 요인을 기준을 구체화하여 명시하였다.

2.3 전문가 검토

1차 비네트의 내용타당도 확보를 위하여 대학기관에서 근무하는 비전임교원과 연구원 5인을 대상으로 전문가 검토를 실시하였다. 전공별로는 교육공학 2인, 교육학 2인, 그리고 영문학 1인을 포함하였으며, 이들은 각각 대학에서 1~6년 간 학습지원 업무를 담당하였다.

Table 2. Expert Panel

	Major	Experience in Higher Education
1	Ph.D. in Educational Technology	1 year
2	Doctoral Candidate in Educational Technology	3 year
3	Doctoral Candidate in Education	6 year
4	Ph.D. in Education	1 year
5	Ph.D. in English Literature	4 year

전문가 검토에서는 선행연구를 통해 도출한 5개 판별요인의 타당도와 각 비네트별로 판별요인 반영, 판단 모호성, 현실성, 교육적 가치 차원에서의 타당도를 검토하였다.

2.4 사용자 테스트

전문가 검토를 통해 수정된 2차 비네트와 질문지는 실제 교육 맥락에서의 적용 가능성 검토를 위하여 교육대학원에 재학 중인 예비교사 22인을 대상으로 사용자 테스트를 진행하였다. 사용자 패널의 구성은 Table 3과 같으며, 사전조사에서 학업과 업무 모두에서 생성형 AI를 활용한 경험이 있다고 응답하였다.

Table 3. User Panel

Participant		N
Semester	1	7
	3	15
Major	Early Childhood Education	7
	Library and Information Science Education	3
	Counseling Studies	2
	Physical Education	3
	Nutrition Education	6
	Business Education	1

사전지식이 없는 학습자의 관점에서 비네트의 이해도와

적절성, 현실성을 검토하기 위하여 전문가 검토에서와 달리 사전에 생성형 AI 표절 판별 요인이나 관련 이슈에 대한 정보는 제공하지 않았다. 또한 응답자에 따라 특정 사례의 맥락이나 생성형 AI 활용 경험에 대한 친숙도 차이로 인해 즉각적인 판단이 어려울 수 있다는 점을 고려하여, 평가를 유보할 수 있는 응답 범주(⑤ 잘 모르겠다)를 포함하였다. 모든 조사가 완료된 후, 이에 대한 교육을 제공하고 토론을 진행하였다.

3. 생성형 AI 활용의 학문적 진실성 판단을 위한 비네트 개발 결과

3.1 전문가 검토 결과

전문가 검토에서는 선행연구를 통해 도출한 AI 활용의 학문적 진실성 판단 요인과 12개의 1차 비네트에 대한 타당도를 확인하였다.

먼저 AI 활용의 학문적 진실성 판단 요인에 대해서는 내용타당도지수(CVI)를 모두 충족하는 것으로 나타났다. 다만 수정 및 보완 의견이 있었던 일부 항목에 대해 의견을 반영하여 수정하였다.

Table 4. Content Validity of Factors for Detecting Misconduct in the Use of Generative AI

Factor	Components	CVI	Decision
AI Contribution	- Proportion of AI-generated content - Proportion of learner modification or editing	1.00	OK
Transparency	- Disclosure of AI use - Awareness by instructor or institution - User intent	1.00	OK
Objective Impartiality	- Alignment with task purpose - Existence of institutional policies	1.00	OK
Authorship	- Attribution of ownership of generated output - Source of ideas - Scope of recognition for creative contribution	1.00	OK
Ethics	- Academic integrity standards - Responsibility for AI-related errors	1.00	OK

검토의견에서는 맥락성과 저작성, 윤리성에 대해서는 모두 동의하였으나, AI 기여도와 관련하여 사례의 맥락이 학습 뿐 아니라 연구와 업무 맥락을 포함하고 있다는 점에서 학습자가 아닌 사용자의 실질적 기여로 수정이 필요하다는 의견을 확인하였다. 또한 투명성과 관련하여 교수자 또는 기관의 인지 여부는 결과에 해당하여 사용자가 부정적 활용에 대한 인식이나 의도가 있었는지 파악하기에 적합하지 않다는 의견과 사용자의 의도는 윤리적 태도에 포함하는 것이 적합하다는 의견을 확인하였다. 이에 따라 AI 기여도는 AI와 사용자의 실질적 기여 수준으로 수정하였으며, 투명성을 AI 사용 명시 여부로 재정의하였다.

12개의 비네트에 대한 전문가 검토 결과는 Table 5와 같다. 다만 전문가 검토가 3점 척도, 5인의 소규모로 진행되었다는 점을 고려하여 기준(≥ 0.78)을 미충족하는 항목에 대해서는 전문가의 개방형 의견을 우선적으로 분석하여 반영하였다[24, 25].

Table5. Content Validity Index of 12 Scenarios

	Clarity	Ambiguity	Realism	Educational Value	mean CVI	Decision
1	1.00	0.40	1.00	1.00	0.85	OK
2	1.00	1.00	1.00	1.00	1.00	OK
3	1.00	1.00	1.00	1.00	1.00	OK
4	1.00	1.00	1.00	1.00	1.00	OK
5	1.00	0.60	1.00	0.60	0.80	OK
6	1.00	0.40	0.60	0.60	0.65	
7	1.00	0.60	1.00	0.60	0.80	OK
8	1.00	0.60	1.00	1.00	0.90	OK
9	0.60	1.00	1.00	1.00	0.90	OK
10	1.00	1.00	1.00	1.00	1.00	OK
11	1.00	0.60	1.00	1.00	0.90	OK
12	1.00	0.40	0.60	0.60	0.65	

검토 의견에서는 전체적으로 학업 상황이 교육학 관련 수업으로 한정되어 있어 보다 다양한 학습 맥락을 포함할 필요가 있다는 의견을 확인하여 수정 반영하였다. 항목별로는 비네트 간 유사성이 높아 활용 목적과 범위 등을 보다 구체적으로 제시하여 맥락성이 명확히 드러날 수 있도록 개선이 필요하다는 의견과 AI의 산출물을 생성받는다는 표현 대신 생성한다고 표현하여 중립성을 확보해야한다는 의견을 확인하였다.

타당도 기준을 미충족한 항목에서는 활용 범위를 구체화하여 다른 사례와 차별성이 필요하다(case 6), 전문성을 요하는 과제로 보기 어렵다(case 12)는 의견을 확인할 수 있었다. 이러한 전문가 의견을 반영하여 Table 6과 같이 수정하였다.

Table 6. Revised Vignettes

	Before	After
4		Added task guidelines: The learner's perspective should be incorporated
5	was generated using AI	generated using AI
6	AI usage rate not specified	Specified that 70% or more of the report was written using AI
10	Received a list from AI	Generated a list using AI
11	Received a list of items from AI	Generated a list of items using AI
12	Field survey	Submission of a policy proposal

3.2 사용자 테스트 결과

개발한 비네트의 적절성과 현실성을 검토하기 위하여 교육대학원에 재학 중인 22명의 대학원생을 대상으로 사용자 테스트를 실시하였다. 각 문항은 5개 응답 범주(1 = 전혀 아니다, 2 = 아니다, 3 = 그렇다, 4 = 매우 그렇다, 5 = 잘 모르겠다)로 구성되었다. 이 중 '잘 모르겠다'는 판단 유보의 의미를 담고 있어 별도 응답 범주로 간주하여 평균 및 표준편차 산출에서는 제외하였다. 이에 따라 1~4점 응답을 대상으로 기술통계를 산출하였으며, 응답 경향을 보다 명확히 파악하기 위하여 긍정 응답 비율(3점 및 4점 응답 비율)을 추가적으로 산출하여 결과를 해석하였다. 그 결과, 교육적 적절성은 평균 3.09점 ~ 3.59점, 현실성은 3.50점 ~ 3.75점으로 나타나 대체적으로 긍정적인 인식을 보인 것으로 나타났다.

Table 7. Results of user testing

	Educational Appropriateness			Realism		
	M	SD	Positive Response	M	SD	Positive Response
1	3.36	0.79	81.8	3.73	0.46	100.0
2	3.32	0.89	72.7	3.64	0.49	100.0
3	3.09	0.87	77.3	3.55	0.51	100.0
4	3.55	0.80	81.8	3.59	0.50	100.0
5	3.50	0.86	77.3	3.73	0.55	95.5
6	3.14	0.77	77.3	3.55	0.60	95.5
7	3.18	1.01	68.2	3.64	0.58	95.5
8	3.50	0.86	77.3	3.73	0.55	95.5
9	3.59	0.80	81.8	3.59	0.50	100.0
10	3.50	1.06	59.1	3.55	0.51	100.0
11	3.36	0.73	86.4	3.50	0.51	100.0
12	3.32	0.78	81.8	3.59	0.59	95.5

긍정 응답 비율을 살펴보면, 적절성에서는 대부분 긍정 응답 비율이 70.0% 이상으로 나타났으나, Case 7(68.2%)과 Case 10(59.1%)은 상대적으로 낮은 수준을 보였다. 다른 Case에 비해 표준편차가 크게 나타났다는 점을 고려하면, 이는 응답자 간 판단 불일치가 크게 나타난 것으로 해석할 수 있다. 이는 해당 사례들이 AI 활용 범위를 명확히 판단하기 어려운 경계적 상황을 반영하고 있기 때문으로 해석된다.

Case 7은 연구자가 연구설계와 본문 작성 전반을 직접 수행한 상태에서 영문 초록 작성을 위해 AI를 활용한 사례로, 이를 언어적 교정을 위한 보조적 활용으로 볼 것인지, 혹은 실질적인 저작 기여로 간주할 것인지에 따라 판단이 달라진 것으로 해석된다.

Case 10은 AI가 제안한 아이디어를 학습자가 구체화한 사례로, AI를 발상 촉진 도구로 인정할 것인지, 혹은 과제의 핵심 창의성을 대체한 것으로 판단할 것인지에 따라 응답이 달라질 수 있다. 특히 과제 평가 기준에 '아이디어의 독창성'이 명시되어 있다는 점에서 독창성의 판단 기준을 아이디어의

출처에 둘 것인지, 학습자의 구체화와 발전 과정에 둘 것인지에 따라 해석이 달라질 수 있는 복합적 사례로 볼 수 있다.

반면에 현실성에서는 12개 비네트 모두 긍정 응답 비율이 95.5% 이상으로 나타나 실제 학습 상황을 충분히 반영하고 있는 것으로 볼 수 있다.

4. 결론 및 제언

본 연구는 생성형 AI 활용 맥락에서 학습자의 학문적 진실성 판단을 측정하고 교육적으로 활용할 수 있는 비네트 기반 도구를 개발하고, 그 타당성을 검토하는 것을 목적으로 수행되었다. 이를 위해 문헌고찰을 통해 AI 활용의 학문적 진실성 판단 요인을 도출하고, 이에 기반하여 초기 비네트를 개발하였으며, 전문가 검토와 사용자 테스트를 통해 최종 비네트와 질문지를 개발하였다.

그 결과 본 연구에서 개발된 AI 활용의 학문적 진실성 판별 비네트가 전반적으로 높은 현실성을 확보한 것으로 나타났다. 이는 비네트의 사례가 학습자들이 생성형 AI를 활용하는 과정에서 실제로 경험 가능한 상황들을 충분히 반영하고 있음을 보여준다. 또한 교육적 적절성 측면에서도 대체로 긍정적인 평가를 받았다. 특히 일부 사례에서 확인된 낮은 교육적 적절성과 높은 응답 분산은 단순한 도구의 한계라기보다 생성형 AI 활용의 학문적 진실성에 대한 윤리적 판단이 본질적으로 맥락 의존적이며 다층적인 특성을 가진다는 점을 반영하는 결과로 해석할 수 있다. 즉, 동일한 상황에서도 학습자의 경험, 신념, AI에 대한 인식에 따라 판단이 달라질 수 있으며, 이러한 특성은 본 도구가 단일한 정답을 측정하기보다는 판단 과정 자체를 드러내는 도구로 기능할 수 있음을 보여준다.

이러한 결과에 따른 본 연구의 시사점은 다음과 같다. 첫째, 생성형 AI 시대의 학문적 진실성 교육은 규범 전달 중심에서 벗어나 상황 기반 판단 교육으로 전환될 필요가 있다. 본 연구의 비네트는 단순한 지식 평가가 아니라, 실제와 유사한 상황에서의 판단을 요구함으로써 학습자가 자신의 기준과 사고 과정을 성찰하도록 유도하는 도구로 활용될 수 있다. 특히 생성형 AI의 활용 과정에서 직면하는 윤리적 문제들이 허용과 금지의 이분법으로 설명하기 어려운 영역이라는 점에서 중요한 교육적 접근이 될 수 있다.

둘째, 본 도구는 AI 활용의 학문적 진실성에 대한 학습자의 인식을 점검하는 진단도구이자 교육용 도구로 활용될 수 있다. 비네트는 학습자의 현재 판단 수준을 진단하는 기능뿐 아니라, 토의, 피드백, 반성적 글쓰기 등과 결합하여 학습자의 윤리적 판단을 확장하는 학습 자료로 기능할 수 있다. 특히 판단이 분산되는 사례는 다양한 관점을 비교하고 논의하는 데 효과적으로 활용될 수 있다는 점에서 교육적 가치가 높다.

다만 본 비네트를 수업자료로 활용할 때에는 교수자의 적절한 교육적 개입이 요구된다. 도덕 발달 모형에 따르면, 윤

리적 판단 역량은 규범의 인지를 넘어 상황 해석과 판단, 성찰의 과정을 통해 형성된다[17]. 이러한 특성을 고려하면, 교수자는 학습자가 사례에 대해 숙고하고 서로 다른 관점을 비교하며 판단 근거를 명료화 하여 재성찰할 수 있도록 지원할 필요가 있다. 이 과정에서 교수자는 전체 학습자의 응답 분포와 자신의 응답을 비교할 수 있는 데이터를 제공하여 학습자가 스스로 자신의 판단 기준을 점검하고 재구성하는 메타인지적 성찰을 경험하도록 할 수 있다. 그리고 이 과정을 통해 학습자는 보다 정교한 AI 사용의 학문적 진실성 판단 기준을 형성할 수 있을 것으로 기대된다.

다음으로 본 연구의 한계와 후속 연구를 위한 제언은 다음과 같다.

첫째, 본 연구는 교육대학원생을 대상으로 사용자 테스트를 실시하였다는 점에서 표본의 제한을 가진다. 생성형 AI 활용에 대한 학문적 진실성 판단은 학습자의 전공 특성, 학업 경험, 생성형 AI 활용 경험 등에 따라 다르게 나타날 가능성이 있으므로, 향후 연구에서는 다양한 전공과 학습 수준을 포함한 집단을 대상으로 도구의 적용 가능성을 검토할 필요가 있다.

둘째, 본 연구는 비네트 기반 질문지의 개발과 타당화를 목적으로 내용 타당성과 사용자 반응 특성을 중심으로 검토하였다. 향후 연구에서는 사례별 반응 양상과 판단 기준의 구조를 보다 정교하게 분석함으로써 생성형 AI 활용에 대한 학습자의 학문적 진실성 판단 유형과 그 특성을 심층적으로 탐색할 필요가 있다.

셋째, AI의 윤리적 활용에 관한 교육에서는 학습자의 판단 결과를 측정하는 것보다 판단 과정을 이해하고 변화를 유도하는 것이 더욱 중요하다. 이러한 측면에서 교육 자료로서 비네트의 활용 가능성을 확인하였으나, 실제 교수학습 맥락에서의 활용 효과는 검증하지 못하였다. 따라서 향후 연구에서는 대학 수업 현장에 본 도구를 적용하여 학습자의 윤리적 성찰, 판단 기준의 변화, 메타인지적 성찰 촉진 효과 등을 실증적으로 검토할 필요가 있다.

이처럼 본 연구는 생성형 AI 활용 맥락에서 학습자의 학문적 진실성 판단을 탐색할 수 있는 비네트 기반 도구를 개발하고 그 타당성을 확보하였다는 점에서 의의를 가진다. 특히 본 도구는 평가를 넘어 학습자의 사고를 드러내고 확장하는 교육적 도구로 활용될 수 있으며, 생성형 AI 시대의 학문적 진실성 교육을 위한 실천적 기반을 제공한다는 점에서 중요한 의미를 가진다.

참고문헌

- [1] International Center for Academic Integrity. (2021). *The fundamental values of academic integrity* (3rd ed.).
- [2] Perkins, M. (2023). Academic integrity considerations of AI large language models in the post-pandemic era: ChatGPT and beyond. *Journal of University Teaching and Learning Practice*, 20(2), 1–24. <https://search.informit.org/doi/10.3316/informit.T2024111300009591751711095>
- [3] Chan, C., & Tsi, L. (2023). The AI revolution in education: Will AI replace or assist teachers in higher education? *Studies in Educational Evaluation*, 83. <https://doi.org/10.48550/arXiv.2305.01185>
- [4] Tang, A., Li, K., Kwok, K., Cao, L., Luong, S., & Tam, W. (2024). The importance of transparency: Declaring the use of generative artificial intelligence (AI) in academic writing. *Journal of Nursing Scholarship*, 56, 314–318. <https://doi.org/10.1111/jnu.12938>
- [5] Cyphert, A. (2023). Generative AI, plagiarism, and copyright infringement in legal documents. *Minnesota Journal of Law, Science & Technology*, 25(2), 49–65. <https://doi.org/10.24926/15529541.3896>
- [6] Hosseini, M., Gordijn, B., Kaebnick, G., & Holmes, K. (2025). Disclosing generative AI use for writing assistance should be voluntary. *Research Ethics*, 21(4), 728–735. <https://doi.org/10.1177/17470161251345499>
- [7] Fraser, K., Dawkins, H., & Kiritchenko, S. (2025). Detecting AI-generated text: Factors influencing detectability with current methods. *Journal of Artificial Intelligence Research*, 82, 2233–2278. <https://doi.org/10.1613/jair.1.16665>
- [8] Chan, C., & Hu, W. (2023). Students' voices on generative AI: Perceptions, benefits, and challenges in higher education. *International Journal of Educational Technology in Higher Education*, 20(1), 43. <https://doi.org/10.1186/s41239-023-00411-8>
- [9] Stokel-Walker, C., & Van Noorden, R. (2023). What ChatGPT and generative AI mean for science. *Nature*, 614, 214–216. <https://doi.org/10.1038/d41586-023-00340-6>
- [10] Lund, B. (2025). AI and academic integrity: Exploring student perceptions and implications for higher education. *Journal of Academic Ethics*, 23, 1545–1565. <https://doi.org/10.1007/s10805-025-09613-3>
- [11] Chan, C., & Colloton, T. (2024). *Generative AI in higher education: The ChatGPT effect*. Routledge.
- [12] Cotton, D., Cotton, P., & Shipway, J. (2023). Chatting and cheating: Ensuring academic integrity in the era of ChatGPT. *Innovations in Education and Teaching International*, 61(2), 228–239. <https://doi.org/10.1080/147>

03297.2023.2190148

- [13] Nature. (2023). Tools such as ChatGPT threaten transparent science; here are our ground rules for their use. *Nature*, 613. <https://doi.org/10.1038/d41586-023-00191-1>
- [14] Perkins, M., Roe, J., Postma, D., McGaughan, J., & Hickerson, D. (2023). Detection of GPT-4 generated text in higher education. *Journal of Academic Ethics*, 22, 89–113. <https://doi.org/10.1007/s10805-023-09492-6>
- [15] Baidoo-Anu, D., & Owusu Ansah, L. (2023). Education in the era of generative artificial intelligence (AI): Understanding the potential benefits of ChatGPT in promoting teaching and learning. *Journal of AI*, 7(1), 52–62. <https://doi.org/10.61969/jai.1337500>
- [16] Lund, B., Wang, T., Mannuru, N., Nie, B., Shimray, S., & Wang, Z. (2023). ChatGPT and a new academic reality: Artificial intelligence-written research papers and the ethics of the large language models in scholarly publishing. *Journal of the Association for Information Science and Technology*, 74(5), 570–581. <https://doi.org/10.1002/asi.24750>
- [17] Rest, J. (1986). *Moral development: Advances in research and theory*. Praeger.
- [18] Hughes, R. (1998). Considering the vignette technique and its application to a study of drug injecting and HIV risk and safer behaviour. *Sociology of Health & Illness*, 20(3), 381–400. <https://doi.org/10.1111/1467-9566.00107>
- [19] Fishbein, M., & Ajzen, I. (2011). Predicting and changing behavior: *The reasoned action approach*. Taylor & Francis.
- [20] Hughes, R., & Huby, M. (2002). The application of vignettes in social and nursing research. *Journal of Advanced Nursing*, 37(4), 382–386. <https://doi.org/10.1046/j.1365-2648.2002.02100.x>
- [21] Jonassen, D. (1997). Instructional design models for well-structured and ill-structured problem-solving learning outcomes. *Educational Technology Research and Development*, 45(1), 65–94. <https://doi.org/10.1007/BF02299613>
- [22] Jenkins, N., Bloor, M., Fischer, J., Berney, L., & Neale, J. (2010). Putting it in context: The use of vignettes in qualitative interviewing. *Qualitative Research*, 10(2), 175–198. <https://doi.org/10.1177/1468794109356737>
- [23] Stokel-Walker, C. (2022, December 9). *AI bot ChatGPT writes smart essays—should professors worry?* *Nature*. <https://doi.org/10.1038/d41586-022-04397-7>
- [24] Lynn, M. (1986). Determination and quantification of content validity. *Nursing Research*, 35(6), 382–385. <https://doi.org/10.1097/00006199-198611000-00017>
- [25] Polit, D., & Beck, C. (2006). The content validity index: Are you sure you know what's being reported? *Research in Nursing & Health*, 29(5), 489–497. <https://doi.org/10.1002/nur.20147>



박하나

- 2008년 이화여자대학교 교육공학과(학사)
- 2015년 이화여자대학교 교육공학과(석사)
- 2022년 이화여자대학교 교육공학과(박사)
- 2024년~현재 경기대학교 연구교수

✚ 관심분야 : 테크놀로지 기반 학습 설계, 협력 학습, AI 윤리

✉ j.hanapark@ewhain.net

부록

〈부록 - 표 1〉 최종 비네트

구분	사례
1	민수는 교육심리학 대학원 수업에서 '자기조절학습 전략이 학업성취에 미치는 영향'을 주제로 10페이지 분량의 학기말 리포트를 제출해야 합니다. 과제 지침에는 선행연구를 검토하고 자신의 관점을 포함하여 작성할 것이 요구되어 있으며, AI 사용에 관한 별도 규정은 명시되어 있지 않습니다. 민수는 제출 마감일 하루 전, ChatGPT에게 "자기조절학습 전략이 학업성취에 미치는 영향에 관한 10페이지 분량의 학술 리포트를 작성해줘. 선행연구를 인용하고 결론도 포함해줘"라고 요청하였습니다. ChatGPT는 금세 리포트를 생성하였고, 민수는 별도의 검토나 수정 없이 표지의 이름과 제출일만 바꾼 뒤 그대로 제출하였습니다.
2	지혜는 대학원 수업에서 '학습 분석(Learning Analytics) 동향 분석'을 주제로 8페이지 분량의 리포트를 제출해야 합니다. 이를 위해 Claude에게 "학습 분석의 최신 동향을 정리한 학술 리포트를 작성해줘"라고 요청하여 완성된 초안을 받았습니다. 이후 AI가 생성한 텍스트에서 단어 일부를 동의어로 교체하고(예: 활용 → 사용, 분석하다 → 검토하다) 문단 순서를 일부 바꾸어 제출하였습니다. 수업 지침과 학과 지침을 모두 확인했지만, AI 활용에 대한 지침이나 별도의 언급은 확인할 수 없어 AI 사용 여부는 표기하지 않았습니다.
3	영호는 교육과정 대학원 수업에서 '2022 개정 교육과정의 방향과 쟁점'에 대한 비평 에세이를 제출해야 합니다. 과제 지침에는 "단순한 내용 요약이 아닌, 자신의 관점을 논거와 함께 비판적으로 서술할 것"이 명시되어 있습니다. 영호는 AI에게 2022 개정 교육과정의 배경, 주요 변화, 쟁점을 정리해 달라고 요청하여 교육과정 개정의 역사와 배경 설명을 생성받아 에세이의 약 60%에 해당하는 부분을 완성하였습니다. 그리고 여기에 자신의 비판적 분석과 의견을 작성하여 최종 과제를 완성하고 제출하였습니다. AI 사용 여부는 어디에도 표기하지 않았습니다.
4	수진이는 연구방법론 수업에서 '교육 분야 AI 연구 동향'을 주제로 문헌 검토 보고서를 제출해야 합니다. 수진이는 Perplexity AI를 활용하여 관련 논문을 찾고, 요약과 분석하도록 요청했습니다. AI를 통해 약 50% 정도를 작성하였고, 나머지는 자신이 직접 작성하여 보고서를 완료하였습니다. 과제 안내에는 학습자의 분석적 관점을 반영하도록 명시되어 있었으나, AI 사용에 관한 별도 지침은 없었습니다. 수진이는 서두에 "본 보고서의 일부는 AI 도구를 활용하여 작성하였음"이라는 한 줄 주석을 달았습니다. 그러나 어떤 도구를 어느 범위에서 어떻게 활용하였는지는 구체적으로 기술하지 않았습니다.
5	준호는 데이터 분석 수업에서 R을 사용하여 학습자 성취 데이터를 분석하고 시각화하는 과제를 제출해야 합니다. 과제 안내문에는 "직접 코드를 작성하고, 각 분석 단계의 의미를 설명하는 보고서를 함께 제출할 것"이라고 명시되어 있습니다. 준호는 GitHub Copilot을 사용하여 데이터 불러오기, 전처리, 회귀분석, 시각화에 필요한 코드 전체를 자동 생성했습니다. 코드가 오류 없이 실행되는 것을 확인한 후 결과물을 그대로 제출하였으며, 보고서의 분석 설명도 Copilot이 생성한 주석을 그대로 옮겨 적었습니다. AI 사용 사실은 코드 어디에도 표기하지 않았습니다.
6	마케팅팀 소속의 현주는 신제품 출시를 앞두고 경쟁사 동향 및 시장 환경을 분석한 보고서를 경영진에게 제출해야 합니다. 팀 내에서는 담당자가 직접 시장을 조사하고 분석하여 작성하는 것이 관행이며, 회사 차원의 AI 사용 관련 내부 정책은 마련되어 있지 않습니다. 현주는 시간이 부족하여 ChatGPT에게 "국내 소비재 시장의 최근 동향과 주요 경쟁사 현황을 분석해줘"라고 요청하여 생성된 내용(약 70%)을 보고서의 핵심 섹션으로 사용하고, 자사 판매 데이터와 간단한 의견을 추가하여 제출하였습니다. AI가 제시한 시장 점유율 수치와 경쟁사 정보의 사실 여부는 별도로 확인하지 않았으며, AI 사용 사실에 대해서도 언급하지 않았습니다.
7	태희는 철학박사과정생으로, 3년에 걸쳐 직접 설계하고 수행한 혼합 연구 결과를 국제 학술지에 투고하려 합니다. 연구 설계, 데이터의 수집과 분석, 논문 본문 전체 작성 모두 본인이 직접 하였습니다. 그러나 영문 초록 작성에서는 어려움을 겪었고, 고민 끝에 자신이 작성한 국문 초록을 AI에 입력하고 이 내용을 바탕으로 국제 학술지 투고용 영문 초록을 250단어로 작성해 달라고 요청하였습니다. AI가 생성한 영문 초록에서 일부 전문용어만 수정한 뒤 그대로 사용하였으며, 논문 어디에도 AI 사용 사실을 표기하지 않았습니다. 투고 예정 학술지의 게재 안내문에는 AI 사용에 관한 정책은 명시되어 있지 않았습니다.
8	소영이는 '2019~2024년 연구동향'을 정리하여 10페이지 분량의 보고서를 제출해야 합니다. 교수자는 "최소 15편의 논문을 직접 읽고 비판적으로 종합하여 작성할 것"을 명시하였습니다. 소영이는 15편의 논문 PDF를 AI에 업로드하고 각 논문의 연구 목적, 방법, 결과, 한계를 3~4줄로 요약해 달라고 요청하였습니다. 그리고 AI가 생성한 요약들을 주제별로 묶고 연결 문장을 추가하는 방식으로 보고서를 완성하였습니다. 논문의 원문을 직접 읽거나, 요약된 내용이 정확한지 확인은 하지 않았습니다. 참고 문헌에는 15편의 논문을 모두 표기하였으나, AI 사용 사실은 밝히지 않았습니다.
9	정민이는 행정학 수업의 팀 프로젝트에서 팀원 5명이 각자 섹션을 나누어 도시 혁신 정책 분석 보고서를 작성하기로 했습니다. 팀 회의에서 역할 분담은 논의하였으나, AI 사용에 대한 합의는 별도로 하지 않았습니다. 정민이는 자신이 담당한 '혁신도시 성과 분석' 섹션 전체를 AI로 생성하여 팀 공유 문서에 업로드하였습니다. 나머지 팀원들은 각자 섹션을 직접 작성하였으며, 정민이가 AI를 사용한 사실은 아무도 알지 못했습니다. 5명 전원의 이름으로 제출되었으며, AI 사용 사실은 보고서에 표기되지 않았습니다.
10	동현이는 미디어 콘텐츠 개발 수업의 과제로 초등학교 대상 교육 게임 기획안을 제출해야 합니다. 과제 평가 기준에는 아이디어의 독창성과 교육적 효용성이 명시되어 있으며, 교수자는 수업 중 "잠시한 아이디어를 스스로 발굴해볼 것"을 강조하였습니다. 동현이는 AI에게 "초등학교 3~4학년 수학 개념 학습을 위한 교육 게임 아이디어 10가지를 구체적으로 제시해줘"라고 요청하여 아이디어 목록을 작성했습니다. 그 중 '분수 개념을 활용한 요리 게임' 아이디어를 선택하여 게임 규칙, 학습 목표, 화면 구성 등을 자신이 직접 구체화한 기획안을 제출하였습니다. 아이디어의 출처가 AI임은 기획안에 밝히지 않았습니다.
11	보람이는 연구방법론 수업의 과제로 설문지를 직접 개발하고 내용 타당도를 검증해서 제출해야 합니다. 과제 안내문에는 "선행 연구를 바탕으로 구성 개념을 정의하고, 연구자가 직접 문항을 개발할 것"이 명시되어 있습니다. 보람이는 AI에게 "중장년층의 AI 리터러시를 측정하는 Likert 5점 척도 설문 문항 20개를 하위 요인별로 개발해줘"라고 요청하여 문항 목록을 생성하였습니다. 이 중 3개 문항의 표현만 수정하고 나머지는 그대로 사용하여 보고서에 '연구자 개발 도구'로 제시하였습니다. 선행연구 기반의 구성 개념 정의 과정과 타당도 절차는 생략하였으며, AI 사용 사실도 표기하지 않았습니다.
12	중학교 교사인 재훈은 교육청으로부터 교내 생성형 AI 활용 지침을 담은 정책 제안서를 제출하도록 요청받았습니다. 제안서는 학교 현장의 전문가로서 재훈의 경험과 판단을 바탕으로 작성하는 것으로 기대되며, 학교 및 교육청 차원의 AI 사용 관련 정책은 아직 마련되어 있지 않습니다. 재훈이는 AI에게 "중학교에서 생성형 AI 사용 지침을 만든다면 포함되어야 할 내용과 구체적인 정책 방안을 제안해줘"라고 요청하여 얻은 답변을 토대로, 학교명과 일부 표현만 수정하여 정책 제안서를 완성하였습니다. 제안서에 AI 활용 사실은 명시하지 않았으며, AI가 제안한 내용이 해당 학교 현장에 실제로 적합한지에 대한 별도 검토는 이루어지지 않았습니다.

Q1-1. 이 사례가 표절, 학문적 윤리 위반에 해당한다고 생각하나요?

- ① 전혀 아니다 ② 아니다 ③ 그렇다 ④ 매우 그렇다 ⑤ 잘 모르겠다

Q1-2. 그렇게 판단한 근거가 무엇인지 해당하는 것을 모두 선택하세요.

- ① AI 사용 여부 표기 및 출처 명시
② 학습자의 실질적 기여 수준(아이디어, 판단, 편집 등)
③ 과제/업무의 특성 또는 유형
④ 기만 또는 속임의 의도성
⑤ 결과물에 대한 소유와 책임
⑥ 기타()

Q1-3. 이 사례에서 AI의 활용이 정당하다고 생각하나요?

- ① 전혀 아니다 ② 아니다 ③ 그렇다 ④ 매우 그렇다 ⑤ 잘 모르겠다

Q1-4. 이 사례에서 AI 사용여부에 대한 고지가 필요하다고 생각하나요?

- ① 불필요하다(학습보조 도구로 사용한 것으로 간주할 수 있다)
② 부분적으로 필요하다
③ 반드시 출처, AI 사용에 대한 표기가 필요하다

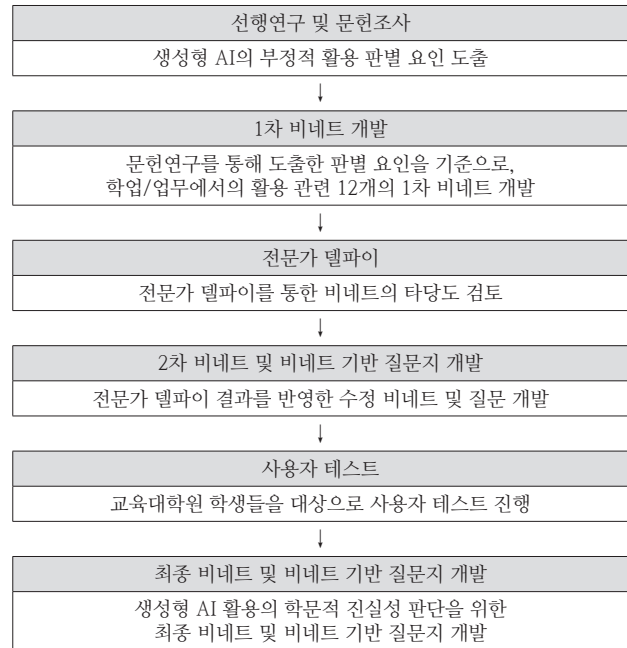
Q1-5. 이 사례에서 결과물에 대한 책임은 학생에게 있다고 생각하나요?

- ① 전혀 아니다 ② 아니다 ③ 그렇다 ④ 매우 그렇다 ⑤ 잘 모르겠다

〈부록 - 그림 1〉 최종 비네트 기반 질문지

〈표 4〉 생성형 AI 활용의 학문적 진실성 판별 요인 내용 타당도

요인	구성요소	CVI	Decision
AI 기여도	- AI 생성 비율 - 사용자의 수정 또는 편집 수준	1.00	OK
투명성	- AI 사용 명시 여부 - 교수자 또는 기관의 인지 여부 - 사용자의 의도	1.00	OK
맥락성	- 과제(업무)의 목적 부합성 - 기관의 정책 존재 여부	1.00	OK
저작성	- 산출물의 저작권 귀속 - 아이디어 출처 - 창작 기여 인정 범위	1.00	OK
윤리성	- 학문적 정직성 규범 - AI 오류 책임 귀속	1.00	OK



〈그림 1〉 연구절차

〈표 6〉 비네트 수정안

	수정 전	수정 후
4		과제 지점 추가(과제 안내에는 학습자의 분석적 관점을 반영하도록 명시되어 있었으나)
5	자동 생성받았습니다.	자동 생성했습니다.
6	AI 활용 비율 명시하지 않음	보고서의 70% 이상이 AI로 작성되었음을 명시함
10	목록을 받았습니다.	목록을 작성했습니다.
11	문항 목록을 생성받았습니다.	문항 목록을 생성하였습니다.
12	현장 조사	정책 제안서 작성으로 수정하여 전문성에 대한 요구 강화

〈표 1〉 비네트 설계안

	과제 맥락	AI 활용 범위			윤리적 쟁점
		AI 활용 범위	학습자 기여	AI 사용 고지	
1	학업 (보고서 작성)	완전 위임	없음	미고지	AI 생성물 무수정 제출: 완전 표절의 기준점
2	학업 (보고서 작성)	완전 위임	최소 (일부 용어 및 문단 배치 수정)	미고지	동의를 교제, 단락 재배열 등 최소한 개입의 효과
3	학업 (에세이 작성)	부분 위임 (60%)	40% 직접	미고지	혼합 기여 상황에서 사용자의 개입 비중
4	학업 (문헌고찰)	부분 위임 (50%) + 부분 고지	50% 직접	부분 고지	혼합 기여 상황에서 사용자의 개입 비중, 고지 방식의 충분성에 따른 효과
5	학업 (코딩)	완전 위임	없음	미고지	코딩 과제 AI 자동 생성: 직접 작성 명시 위반
6	직무 (시장 분석)	부분 위임 (70%)	30%	미고지	직무 맥락 AI 활용: 사실 미검증 및 사용 여부 미고지
7	연구 (영문 초록 작성)	완전 위임 부분 위임 (70%)	본문 전체 직접	미고지	언어 도구 활용의 윤리적 허용 경계
8	학업 (문헌고찰)	보조(요약)	구조화 직접	미고지	논문 직독 없이 AI 요약 기반 보고서 작성
9	학업 (팀 프로젝트)	부분 (팀원 1인만 활용)	팀원 1인 외 직접	미고지 (팀 내 미고지)	집단 책임 귀인과 팀 내 기만
10	학업 (기획안 개발)	부분 (아이디어 도출)	구체화 직접	미고지	창의성 평가 과제에서 AI 아이디어 원천
11	학업 (도구 개발)	문항 생성 전량	3개 아이템 수정	허위 표기	'연구자 개발' 허위 표기: 연구 도구 조작
12	직무 (정책 제안)	전면 대체+최소 수정	최소	미고지	전문가 경험 및 판단의 AI 대체