

# 대학 연구자를 위한 생성형 AI 연구윤리 가이드



교육부



한국연구재단



KUCRE  
대학연구윤리협의회

# 대학 연구자를 위한 생성형 AI 연구윤리 가이드

※ 본 가이드는 연구의 전(全) 과정에서 생성형 AI 활용과 관련하여 발생할 수 있는 연구윤리 문제를 예방하고 연구자의 책임 있는 연구 수행을 지원하기 위해 작성되었습니다. 각 대학의 연구 환경에 맞는 생성형 AI 활용 가이드 수립 및 운영에 참고하시기 바랍니다.



교육부



한국연구재단



KUCRE  
대학연구윤리협의회

## 목적 및 활용 01

### 제1장 03

#### 연구에서 활용되는 생성형 AI

### 제2장 07

#### 연구 설계 단계에서의 생성형 AI 활용

|             |    |
|-------------|----|
| 1) 선행 연구 검토 | 9  |
| 2) 연구 문제 설정 | 12 |
| 3) 가설 생성    | 14 |
| 4) 연구 방법 결정 | 17 |

### 제3장 19

#### 연구 수행 단계에서의 생성형 AI 활용

|           |    |
|-----------|----|
| 1) 데이터 수집 | 21 |
| 2) 데이터 분석 | 23 |
| 3) 결과 해석  | 27 |

### 제4장 31

#### 연구 보고 단계에서의 생성형 AI 활용

|            |    |
|------------|----|
| 1) 논문 작성   | 33 |
| 2) 참고문헌 정리 | 36 |
| 3) 논문 검토   | 40 |

### 제5장 43

#### 생성형 AI 활용 후 인용

|                     |    |
|---------------------|----|
| 1) 생성형 AI 활용 사실의 공개 | 44 |
| 2) 저널별 정책 확인하기      | 47 |

### 제6장 53

#### 연구자를 위한 생성형 AI 활용 체크리스트

### 부록 57

#### 생성형 AI 활용을 위한 책임 있는 연구 수행 10가지 원칙

### 참고문헌 59

# 대학 연구자를 위한 생성형 AI 연구윤리 가이드

---

## 목적 및 활용

- 최근 생성형 AI 기술이 빠르게 확산되면서 연구 전반에 걸쳐 다양한 방식으로 활용되고 있음. 문헌 검색, 번역, 논문 작성 지원, 데이터 분석 등 연구의 여러 단계에서 생성형 AI 활용이 증가하고 있으며, 이에 따라 새로운 연구윤리 이슈도 함께 제기되고 있음.
- 본 가이드는 생성형 AI 활용 과정에서 발생할 수 있는 연구윤리 문제를 예방하고, 연구자(교원, 연구원, 대학원생 등)가 책임 있는 방식으로 AI를 활용할 수 있도록 안내하기 위해 마련되었음. 특히 실제 연구 과정에서 발생할 수 있는 사례를 중심으로 생성형 AI 활용의 위험 요인과 대응 방안을 제시하고자 함.
- 본 가이드는 생성형 AI의 기술적 활용 방법을 설명하는 데 목적이 있는 것이 아니라, 연구자가 연구의 전 과정에서 생성형 AI를 책임 있게 활용하기 위한 윤리적 기준을 제시하는 데 목적이 있음. 이를 통해 연구의 투명성, 재현성, 신뢰성을 유지하고 연구 진실성을 확보하는 데 도움을 주고자 함.
- 본 가이드는 다음과 같은 내용으로 구성되어 있음.
  - 1장에서는 생성형 AI란 무엇이며, 어떻게 연구에 활용되고 있는지 설명함.
  - 2장에서는 연구 설계 단계에서 생성형 AI 활용 시 발생할 수 있는 문제와 유의 사항을 다루며, 선행 연구 검토, 연구 문제 설정, 가설 생성 등에서 고려해야 할 사항을 제시함.
  - 3장에서는 연구 수행 단계에서 생성형 AI 활용과 관련된 주요 윤리 이슈를 설명하며, 데이터 수집, 데이터 분석, 결과 해석 과정에서 발생할 수 있는 문제와 예방책을 다룸.
  - 4장에서는 연구 보고 단계에서 생성형 AI 활용과 관련된 문제를 설명하며, 논문 작성, 참고문헌 정리, 논문 검토 과정에서 발생할 수 있는 연구윤리 문제와 대응책을 제시함.
  - 5장에서는 생성형 AI 활용 시 필요한 인용 및 공개 방법과 주요 학술지의 생성형 AI 정책을 소개함.
  - 6장에서는 연구자가 생성형 AI를 활용할 때 점검해야 할 주요 사항을 체크리스트 형태로 제시함.
- 본 가이드가 연구를 수행하는 모든 사람이 생성형 AI 활용 과정에서 연구윤리를 올바르게 이해하고, 책임 있는 연구를 실천하는 데 실질적인 도움이 되기를 기대함.

# 대학 연구자를 위한 생성형 AI 연구윤리 가이드



연구에서 활용되는  
생성형 AI

# 제1장

## 제1장

# 연구에서 활용되는 생성형 AI

● 생성형 AI(Generative Artificial Intelligence)는 대규모 데이터를 학습하여 텍스트, 이미지, 코드, 음성 등 새로운 콘텐츠를 생성할 수 있는 기술을 의미함.<sup>1)</sup>

- 최근 대규모 언어모델(Large Language Models, LLMs)의 발전으로 생성형 AI는 연구 활동 전반에서 활용되기 시작하였으며, 문헌 검토, 데이터 분석 지원, 논문 작성 보조,<sup>2)</sup> 번역 등 다양한 연구 활동에 활용되고 있음.

● 연구 과정에는 다음과 같은 다양한 유형의 생성형 AI 도구가 사용되고 있음.

- ① 텍스트 생성형 AI: 문헌 요약, 번역, 학술적 문장 교정과 감수(proofreading), 윤문 작업, 초안 작성, 코드 생성 등을 지원(예시: ChatGPT, Gemini, Claude)
- ② 문헌 탐색 및 연구 지원 AI: 연구 논문 탐색 및 문헌 분석을 지원(예시: Elicit, Scite, Semantic Scholar AI)
- ③ 코드 및 데이터 분석 지원 AI: 연구데이터 분석이나 코드 작성을 지원(예시: GitHub Copilot, Code Interpreter)
- ④ 이미지 생성형 AI: 연구 발표 자료나 시각 자료 제작에 활용(예시: DALL·E, Midjourney)

● 생성형 AI는 연구 생산성과 효율성을 향상시키는 유용한 도구이지만, 그 활용 과정에서 다음과 같은 문제가 발생할 수 있음.

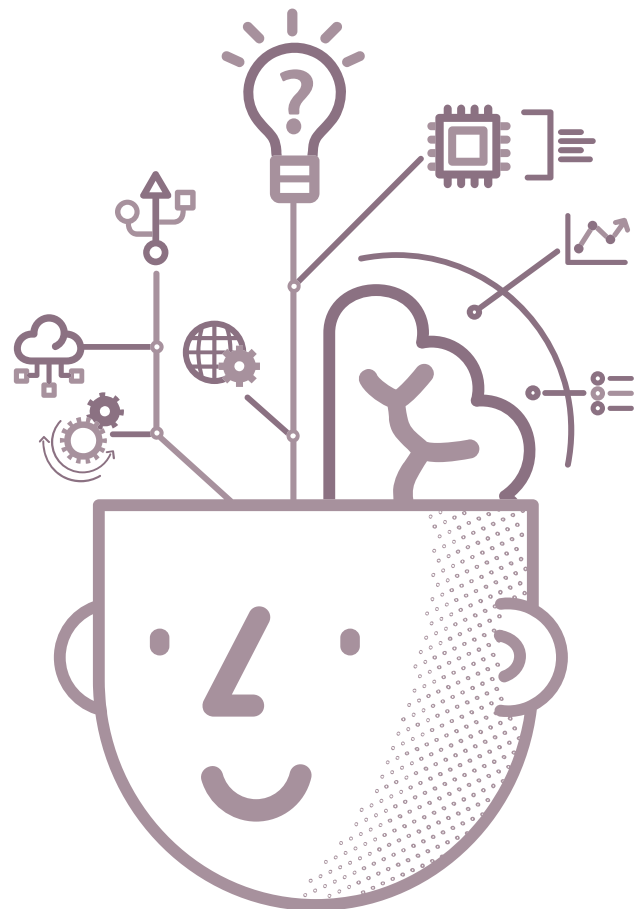
- ① 연구 진실성 문제: 생성형 AI가 응답한 내용을 검증 없이 활용하면 연구 결과의 객관성, 신뢰성, 재현성 등이 훼손될 가능성이 있음.
  - 생성형 AI는 기존의 다양한 문헌과 데이터를 학습하여 콘텐츠를 생성하기 때문에 제대로 된 인용 표시 없이 사용하는 경우, 부적절한 인용 등의 연구윤리 문제가 발생할 수 있음.
- ② 정보의 정확성 문제: 생성형 AI는 사실과 다른 정보를 생성하거나 존재하지 않는 문헌을 인용하는 등 환각(Hallucination) 현상을 일으킬 수 있음.
- ③ 저작권 문제: 생성형 AI는 인터넷상의 방대한 텍스트, 이미지, 논문 등을 학습하여 개발되므로, 저자(저작권자)의 허락 없이 저작물이 학습 데이터로 활용되었을 가능성이 있으며, 그 결과 기존 저작물과 유사한 내용이 생성될 가능성이 있음.

1) Bommasani et al., 2021; Dwivedi et al., 2023.

2) van Dis et al., 2023.

④ 연구 보안 및 개인정보 문제: 연구 데이터나 미공개 연구 결과를 생성형 AI에 입력하면 연구 데이터 유출이나 개인정보 침해가 발생할 수 있음.<sup>3)</sup>

⑤ 연구 책임성 문제: 생성형 AI를 논문 작성에 활용했다고 하더라도, 연구의 책임 주체가 될 수 없으며, 연구 결과에 대한 최종 책임은 연구자에게 있음.<sup>4)</sup>



3) European Commission, 2025.

4) ICMJE, 2023.

# 대학 연구자를 위한 생성형 AI 연구윤리 가이드

---

# 연구 설계 단계에서의 생성형 AI 활용

- 1) 선행 연구 검토
- 2) 연구 문제 설정
- 3) 가설 생성
- 4) 연구 방법 결정

## 제2장

## 제2장

### 연구 설계 단계에서의 생성형 AI 활용

- 연구 설계 단계는 연구자가 연구를 수행하기 전에 연구의 목적, 연구 문제, 가설, 연구 방법 등을 체계적으로 계획하는 과정을 의미함. 이 단계에서는 연구의 방향과 범위를 설정하고, 연구 수행을 위한 이론적·방법론적 틀을 마련하게 됨.
    - 연구 설계 단계는 연구의 전체 과정에서 매우 중요한 단계로, 이 단계에서 연구의 목적과 방법이 명확하게 설정되어야 이후 연구 수행과 결과 해석이 체계적으로 이루어질 수 있음. 연구 설계가 명확하지 않으면 연구 결과의 신뢰성과 타당성이 저하될 수 있음.<sup>5)</sup>
  - 일반적으로 연구 설계 단계에서는 다음과 같은 활동이 이루어짐.
    - ① 선행 연구 검토: 기존 연구를 탐색하고 분석하여 연구 주제와 관련된 이론적 배경과 연구 동향을 파악함.
    - ② 연구 문제 설정: 연구를 통해 해결하고자 하는 핵심 질문 또는 연구 문제를 정의함.
    - ③ 가설 설정: 연구 문제를 검증하기 위해 변수 간의 관계에 대한 가설을 설정함.
    - ④ 연구 방법 결정: 연구 목적에 적합한 연구 방법을 선택하고 연구 설계를 구체화함.
    - ⑤ 연구 대상 및 자료 수집 방법 결정: 연구 대상, 표본, 데이터 수집 방법 등을 결정함.
    - ⑥ 연구 윤리적 검토: 인간이나 동물을 실험 대상으로 하여 생명윤리를 따져야 하는 연구인지 등을 판단함.
- ※ ⑤ 연구 대상 및 자료 수집 방법 결정, ⑥ 연구 윤리적 검토 등은 생성형 AI가 개입할 여지가 적은 영역이기 때문에 본 가이드에 기술하지 않음.

5) Creswell & Creswell, 2018; Saunders et al., 2019.

## 1) 선행 연구 검토

- (정의) 선행 연구 검토는 기존 연구를 탐색하고 분석하여 연구 주제와 관련된 이론적 배경과 연구 동향을 파악하는 과정임. 이를 통해 연구자는 기존 연구의 한계를 확인하고 새로운 연구 문제를 도출할 수 있음.
- (AI 활용) 최근에는 생성형 AI(Elicit, SciSpace, Perplexity 등)를 활용하여 논문 탐색, 문헌 요약, 연구 동향 분석 등을 수행하는 사례가 증가하고 있음. 이러한 도구는 연구자가 방대한 학술 자료를 빠르게 탐색하고 정리하는 데 도움을 줄 수 있음.
- (연구윤리 문제) 생성형 AI를 활용하여 선행 연구를 검토하는 과정에서 다음과 같은 문제가 발생할 수 있음.
  - ① 문헌 접근의 한계: 생성형 AI는 저작권 보호 또는 구독 기반<sup>6)</sup>으로 제공되는 학술 논문이나 전문 서적에 대한 접근이 제한될 수 있음. 따라서 AI가 제시하는 문헌 정보만을 기반으로 선행 연구를 검토할 경우 중요한 선행 연구가 누락될 가능성이 있음.

---

6) 저널은 크게 구독 기반(Subscription-based)과 오픈 액세스(Open Access) 두 가지 방식으로 나뉜다. 구독 기반 저널은 독자가 구독료를 지불해야 논문에 접근할 수 있지만, 오픈 액세스 저널은 저자가 논문 출판 비용(APC: Article Processing Charge)을 부담하고 독자는 무료로 논문에 접근할 수 있다.



이 그림은 Chat GPT에 “한국과 터키의 민주주의 후퇴 비교 연구 논문 작성”을 요청하여 얻은 참고문헌의 여러 가지 문제를 보여줌.

- ① PDF 발췌문: 웹상에 PDF 형식으로 있는 문서를 기반으로 작성하여 그 출처가 불분명하거나, 저작권을 침해하는 경우가 발생함.
- ② 최신 연구 미반영: 최신 연구 결과를 충분히 반영하지 못하거나, 검색 범위 및 데이터 접근성의 한계로 인해 특정 자료에 편중된 결과를 제시할 수 있음.
- ③ 서지 형식의 비완전성: 서지 형식이 틀렸거나 표기가 완전하지 않음.

## 7. 결론

한국과 터키의 민주주의 후퇴는 속도·가시성·견제장치 밀도·시민사회 조직화에서 갈렸다. 급격한 후퇴에 직면한 한국은 제도·사회가 고가시성 충격에 집단적으로 반응하여 법치 절차로 복원했다. 반면 터키의 점진적 권위주의화는 장기적 제도 포획과 시민사회의 분절을 초래하여 회복의 문턱을 높였다. 이 비교는 긴급권의 엄격한 설계, 사법·언론 독립의 제도화, 시민사회 보호가 민주주의 복원력의 핵심임을 보여준다. 향후 연구는 미시적 동원 메커니즘(누가, 언제, 어떤 비용·편익에 반응하는가)과 국제연계(제재·지원·명예규범)의 조건의 효과를 정량적으로 검증해야 한다. Freedom House AP News Journal of Democracy

## 참고문헌(선정)

- Bermeo, N. (2016). On Democratic Backsliding. *Journal of Democracy*, 27(1), 5–19. Journal of Democracy ②번 예시
- Levitsky, S., & Ziblatt, D. (2018). *How Democracies Die*. Crown. (개요 PDF·발췌) ①번 예시  
newuniversityineconomics.org bpb-us-e1.kpmucdn.com
- Esen, B., & Gümüşcü, Ş. (2016). Rising Competitive Authoritarianism in Turkey. *Third World Quarterly*, 37(9), 1581–1606. (오픈 액세스판) Stiftung Wissenschaft und Politik (SWP) ②번 예시
- Freedom House. (2025). *Freedom in the World 2025: South Korea, Turkey*. Freedom House v.1 ③번 예시
- Economist Intelligence Unit. (2025). *Democracy Index 2024* (테이블 2). d1oqtan6gyr07.cloudfront.net ③번 예시
- AP/Reuters/워싱턴포스트 등 현안 보도(2024–2025): 한국의 탄핵·파면·조기대선, 터키의 야권 지도자 체포 및 시위 진압. AP News Reuters The Washington Post v.1 ③번 예시
- IISS *Strategic Comments: South Korea's political crisis* (2025.4). IISS ③번 예시



+ 무엇이든 물어보세요



ChatGPT는 실수를 할 수 있습니다. 중요한 정보는 재차 확인하세요.

출처: “한국과 터키의 민주주의 후퇴 비교 연구 논문 작성”을 명령한 후 나온 서지사항을 캡처하였음.” ChatGPT 5.1, 2025.11.01.

- ② 문헌 정보의 정확성 문제: 생성형 AI는 존재하지 않는 논문이나 부정확한 인용 정보를 생성하는 환각(Hallucination)현상이 발생할 수 있음. 이러한 정보를 검증 없이 활용할 경우 선행연구 검토의 신뢰성이 저하될 수 있음.<sup>7)</sup>

7) Ji et al., 2023.

③ **문헌 편향의 문제**: 생성형 AI는 학습 데이터와 접근가능한 데이터에 따라 특정 연구나 출판물에 편향된 결과를 제시할 수 있음.

- 특히 구독 기반 저널을 인용하지 못하거나, 반대로 게재 철회된 논문을 인용하거나, 신뢰성을 담보하기 어려운 부실 학술지 논문을 인용하여 연구 동향을 왜곡하여 이해할 위험이 있음.

#### AI를 활용한 선행 연구 활용 시 발생할 수 있는 문제점

##### ① 선택적 인용 및 인용 편향(Selective Citation / Citation Bias)

생성형 AI가 특정 연구 방향으로 선행 연구를 제시할 경우 연구자는 연구를 선택적으로 인용하게 될 가능성이 있음.

- 연구자는 선행 연구를 균형있게 검토해야 할 책임이 있음.
- 특정 연구 결과만을 강조하는 선택적 인용이 발생할 경우 연구의 객관성과 학문적 신뢰성이 훼손될 수 있음.

##### ② 연구의 이론적 배경 왜곡 가능성

생성형 AI가 제시한 선행 연구나 이론적 설명이 연구 분야의 실제 논의 흐름과 일치하지 않을 경우 연구의 이론적 배경이나 문제 설정이 지엽적이거나 편향될 가능성이 있음.

- 선행 연구를 제대로 수행하지 않았다고 해서 연구부정행위에 해당할 가능성은 낮지만, 좋은 연구가 아닐 가능성은 높음.

● **(대응방법)** 생성형 AI는 실제 존재하지 않는 문헌을 생성하거나, 서지 정보가 부정확한 참고문헌을 제시할 수 있음. 따라서 연구자는 생성형 AI가 제시한 참고문헌을 그대로 사용하지 않고 반드시 원문을 확인하여 검증하는 과정이 필요함.

① **학술 데이터베이스를 통한 교차 확인**: 생성형 AI가 제시한 논문이나 저널 정보는 Google Scholar, Scopus, Web of Science, PubMed 등 학술 데이터베이스를 통해 실제 존재 여부와 서지 정보를 확인해야 함.

② **DOI 확인**: 논문 제목, 저자, 저널명뿐만 아니라 DOI(Digital Object Identifier)<sup>8)</sup>를 확인하면 문헌의 실제 존재 여부를 보다 정확하게 검증할 수 있음.

③ **원문 직접 확인**: AI가 제공한 요약이나 인용 정보에 의존하지 말고, 해당 논문의 원문을 직접 확인하여 연구 내용과 인용 정보가 정확한지 검토해야 함.

④ **선행 연구 확인**: 생성형 AI를 활용하여 선행연구를 탐색하더라도, 구독 기반 유료 저널과 학술도서 등 주요 학술자료가 충분히 반영되도록 연구자가 직접 확인·검토해야 함.

8) DOI는 디지털 콘텐츠에 부여되는 고유 식별자로서, 논문이나 디지털 문서에 붙이는 영구적인 주소라고 할 수 있다. 일반적인 URL과 달리 시간이 지나거나 논문 URL이 변경되더라도 해당 논문페이지로 이동된다.

## 2) 연구 문제 설정

- (정의) 연구 문제 설정은 연구자가 탐구하거나 해결하고자 하는 핵심 질문을 정의하는 과정임. 연구 문제는 연구의 목적과 방향을 결정하는 핵심 요소로, 전통적인 가설 기반 연구에서는 선행연구 검토를 통해 확인된 연구 공백이나 사회적·학문적 필요성을 바탕으로 도출되며, 데이터 기반 연구에서는 데이터 탐색 및 분석 과정에서 발견된 의미 있는 패턴이나 현상을 통해 형성·발전될 수 있음.
  - 연구자는 연구 문제를 명확하게 설정함으로써 연구의 범위와 목적을 구체화하고, 이후 연구 설계와 분석 과정의 방향을 체계적으로 정립할 수 있음.<sup>9)</sup>
- (AI 활용) 최근에는 생성형 AI를 활용하여 연구 아이디어를 탐색하거나 연구 문제를 구체화하는 사례가 증가하고 있음.
  - 예를 들어 ChatGPT, Gemini, Claude 등의 생성형 AI는 특정 연구 주제와 관련된 연구 질문을 제안하거나 연구 주제를 구체화하는 데 활용될 수 있음. 이러한 도구는 다양한 관점에서 연구 아이디어를 제시함으로써 연구자가 새로운 질문을 탐색하는 데 도움을 줄 수 있음.
- (연구윤리 문제) 생성형 AI를 활용하여 연구 문제를 설정하는 과정에서 다음과 같은 문제가 발생할 수 있음.
  - ① 연구 독창성 저해 가능성: 생성형 AI가 제시하는 연구 아이디어에 의존할 경우, 연구자의 독창적인 문제 설정 능력을 저해하는 문제를 초래할 수 있음.
  - ② 부정확한 정보 기반 연구 문제 설정: 생성형 AI는 학습 데이터의 한계나 환각(Hallucination) 현상으로 인해 부정확한 정보나 존재하지 않는 연구 동향을 제시할 수 있음. 이러한 정보를 검증 없이 활용할 경우 연구 문제 설정의 타당성이 저하될 수 있음.
  - ③ 연구 공백에 대한 오해 가능성: 생성형 AI는 모든 학술 자료를 완전히 반영하지 못하기 때문에 실제로는 이미 충분히 연구된 주제를 새로운 연구 공백으로 제시할 가능성이 있음. 특히 구독 기반 학술지나 최신 연구가 충분히 반영되지 않을 수 있음.

9) Creswell & Creswell, 2018.



= ChatGPT ▾



ESG와 직원 voice behavior 연구에서 아직 연구되지 않은 주제를 알려줘.



정부의 ESG 메시지가 직원의 ESG voice behavior에 미치는 영향에 대한 연구는 아직 충분히 이루어지지 않았으니, 정부 ESG 메시지와 직원의 ESG Voice behavior로 연구 해봐.

하지만, 실제로는 정부의 ESG 메시지와 직원 voice behavior 간의 연구는 이미 다양한 맥락에서 활발히 이루어져 있음. 예를 들면, Detert & Burris (2007), Morrison (2011), Liang et al. (2012) 등

왜 이런 문제가 생기나?

- ① 생성형 AI는 모든 학술 문헌을 반영하지 못함.
- ② 특히 유료 저널(paywalled journals) 연구가 누락될 수 있음.

● (대응방법) 생성형 AI를 활용하여 연구 문제를 설정할 때 발생할 수 있는 문제를 예방하기 위해 연구자는 다음과 같은 사항을 고려해야 함.

- ① AI를 보조 도구로 활용: 생성형 AI는 연구 아이디어 탐색을 위한 보조 도구로 활용하고, 연구 문제의 최종 설정은 연구자가 선행연구 분석과 학문적 판단을 바탕으로 수행해야 함.
- ② 비판적 검토 과정 유지: 연구자는 생성형 AI가 제시하는 연구 질문이나 아이디어를 그대로 사용하는 것이 아니라, 기존 연구와의 차별성, 연구 공백의 타당성, 연구의 학문적 기여 등을 비판적으로 검토해야 함.
- ③ 정보의 정확성 검증: 생성형 AI가 제시하는 연구 동향이나 관련 연구 정보는 Google Scholar, Scopus, Web of Science 등 학술 데이터베이스를 통해 실제 문헌의 존재 여부와 내용을 확인해야 함.
- ④ 충분한 선행연구 검토 수행: 생성형 AI의 결과에만 의존하지 말고, 주요 학술 데이터베이스와 구독 기반 유료 저널, 학술서적 등을 포함하여 선행연구가 충분히 검토되도록 해야 함.

※ 생성형 AI는 연구 아이디어 탐색을 지원하는 도구로 활용할 수 있으나, 연구 문제 설정과 연구 공백의 판단은 반드시 연구자의 선행연구와 학문적 판단에 기반하여 이루어져야 함.

### 3) 가설 생성

- (정의) 가설 생성 단계는 연구자가 연구 문제를 바탕으로 변수 간의 관계를 설명하거나 예측하기 위한 과정임. 가설은 연구에서 검증할 수 있는 명확한 명제로 제시되어야 하며, 이후 연구 설계와 데이터 분석을 통해 검증될 수 있어야 함.
- (AI 활용) 최근에는 생성형 AI를 활용하여 연구 가설을 세우거나 다양한 변수 간 관계에 대한 아이디어를 얻는 사례가 증가하고 있음.
- (연구윤리 문제) 생성형 AI를 활용하여 가설을 생성하는 과정에서 다음과 같은 문제가 발생할 수 있음.
  - ① 검증 불가능한 가설 설정 가능성: 생성형 AI가 제시하는 가설은 실제 연구에서 측정하거나 검증하기 어려운 변수 관계를 포함할 수 있음. 이러한 가설을 그대로 활용할 경우 연구의 실증적 검증 가능성이 낮아질 수 있음.
  - ② 대안 가설의 누락: 연구자가 다양한 설명 가능성을 검토하지 않고 단일 가설에만 집중할 경우 과학적 설명의 타당성이 약화될 수 있음.
  - ③ 생성형 AI가 제시한 가설을 확인 없이 채택: 생성형 AI는 다양한 변수 간 관계를 바탕으로 가설을 제시할 수 있으나, 이러한 가설은 반드시 학술적 근거를 기반으로 생성되는 것은 아님. 연구자가 AI가 제시한 가설을 선행연구 검토나 이론적 검증 없이 그대로 채택할 경우 가설의 학문적 타당성이 확보되지 않을 수 있음.

#### 가설 생성 단계에서의 유의사항: HARKing & Anchoring

- 1) HARKing<sup>10)</sup>은 “Hypothesizing After the Results are Known”의 약자로, 연구 결과를 먼저 확인한 뒤 그 결과에 맞추어 사후에 가설을 설정하는 행위를 의미함.
- 즉, 실제 연구 과정에서는 데이터를 먼저 분석하였음에도, 논문에서는 마치 가설을 먼저 설정하고 이를 검증한 것처럼 서술하는 것임. 일반적으로 과학 연구에서는 이론이나 선행연구를 바탕으로 가설을 설정하고, 이를 경험적 자료를 통해 검증하는 방식이 널리 활용 됨.
- 그러나 HARKing은 데이터를 먼저 수집·분석한 후 그 결과에 부합하도록 가설을 재구성하면서도, 논문에서는 이를 사전에 설정된 가설인 것처럼 제시함.
- 이러한 방식은 실제 연구 과정과 논문에 보고된 연구 과정 사이의 괴리를 초래하며, 연구 결과의 재현 가능성과 과학적 신뢰성을 저해할 수 있다는 점에서 문제적 연구관행(questionable research practice)으로 지적됨.

10) Kerr, N. L., 1998.

2) Anchoring<sup>11)</sup>은 의사결정 과정에서 처음 접한 정보(앵커)가 이후의 판단에 불균형적인 영향을 미치는 인지적 편향으로, 처음 제시된 수치나 정보가 임의적이고 명백히 무관한 경우에도 이후의 판단에 영향을 미치는 현상을 말함.



출처: “대학원생이 생성형 AI에게 연구 질문을 물어보고 대답을 보고 그 대답에 빠져서 다른 생각을 못하는 그림 4컷으로 그려봐”, ChatGPT 5.4, 2026.04.26.

11) Tversky & Kahneman, 1974.

● (대응방법) 생성형 AI를 활용하여 가설을 생성하는 과정에서 발생할 수 있는 문제를 예방하기 위해 연구자는 다음과 같은 사항을 고려해야 함.

① **가설의 검증 가능성 확인:** 연구자는 생성형 AI가 제시한 가설이 실제 연구에서 측정 가능한 변수와 연구 설계를 통해 검증 가능한지를 확인해야 함.

- 특히 변수의 정의, 측정 방법, 데이터 수집 가능성 등을 검토하여 연구자 스스로 실증적 검증이 가능한 가설을 설정해야 함.

② **대안 가설 검토:** 연구자는 단일 가설에만 의존하지 않고 경쟁 가설이나 대안 가설을 함께 고려하여 연구 설계를 구성해야 함.

③ **이론적 근거 및 선행연구 확인:** 생성형 AI가 제시한 가설은 반드시 선행연구와 이론적 논의를 통해 검증해야 함. 연구자는 주요 학술 데이터베이스와 학술 문헌을 확인하여 해당 가설이 기존 연구와 어떤 관계에 있는지 검토하고, 이론적 근거를 기반으로 가설을 설정해야 함.

④ **AI 제안에 대한 인지적 고정 방지:** 생성형 AI가 제시한 가설을 그대로 수용하기보다는 다양한 변수 관계와 설명 가능성을 함께 검토해야 함.

- 생성형 AI는 연구자의 사고를 대체하는 수단이 아니라 보조적 지원 도구로 활용되어야 하며, 연구자가 먼저 충분히 사고하고 내용을 구성한 후 AI에게 질의·검토하는 방식으로 사용하는 것이 바람직함.
- 연구자는 AI가 제시한 가설을 하나의 참고 아이디어로 활용하고, 추가적인 선행연구 검토와 이론적 분석을 통해 가설을 수정·보완할 필요가 있음.

#### 4) 연구 방법 결정

- (정의) 연구 방법 결정 단계는 연구문제와 가설을 검증하기 위해 가장 적절한 연구설계, 자료수집 방식, 표본추출 방법, 변수 측정도구, 분석기법 등을 선택·구체화하는 과정임.
  - 연구방법의 결정은 판단은 연구의 타당성, 신뢰성, 재현가능성에 직접적인 영향을 미침.
- (AI 활용) 생성형 AI는 연구주제에 적합한 정량적 또는 질적 연구방법을 제안하거나, 설문조사, 실험, 인터뷰, 사례연구 등 연구설계를 제안해 줌.
- (연구윤리 문제) 생성형 AI를 활용하여 연구방법을 결정하는 과정에서 다음과 같은 문제가 발생할 수 있음.
  - ① 부적절한 연구방법 제시: AI가 연구 맥락이나 환경을 충분히 이해하지 못한 채 부적절한 방법이나 존재하지 않는 방법을 제안할 수 있음.
    - 부적절한 방법 제시: 방법 자체가 연구수행에 부적절하거나 존재하지 않는 방법을 제시하는 경우
    - 현실 가능성이 결여된 방법 제시: 연구실에 있는 실험 장비, 인력 등 연구환경을 고려하지 못하기 때문에 수행할 수 없는 연구 방법을 제시하는 경우
  - ② 방법의 기계적 채택: 연구자가 방법의 전제조건이나 한계를 검토하지 않고 AI 추천을 그대로 사용할 경우 연구의 재현성과 타당성을 담보할 수 없음.
  - ③ 학문 분야 특수성 무시: 분야별 연구 관행, 윤리기준, 자료 특성을 반영하지 못한 일반적 방법을 제시할 수 있음.
- (대응방법)
  - ① 전문가 검토: 지도교수, 공동연구자, 통계 전문가, IRB 담당자 등의 검토를 거쳐야 함.
  - ② 재현 가능한 연구설계: 연구 절차, 분석 계획 등을 상세히 기록하여 결과 재현이 가능하도록 노력해야 함.
  - ③ AI를 보조 도구로 활용: 최종 연구 방법은 연구자가 직접 결정해야 하며, 그 책임도 연구자에게 있음.

# 대학 연구자를 위한 생성형 AI 연구윤리 가이드

---

# 연구 수행 단계에서의 생성형 AI 활용

- 1) 데이터 수집
- 2) 데이터 분석
- 3) 결과 해석

## 제3장

## 제3장

### 연구 수행 단계에서의 생성형 AI 활용

- 연구 수행 단계는 연구 설계 단계에서 수립한 계획에 따라 실제로 연구를 진행하는 과정으로, 연구자는 이 단계에서 데이터를 수집하고 분석하며 연구 결과를 도출하게 됨.
  - 이 과정에서는 연구 설계에서 설정한 연구 질문과 가설을 검증하기 위해 체계적인 자료 수집, 실험과 조사가 이루어져야 하며, 그 과정에서 연구의 객관성·정확성·재현성을 확보하는 것이 중요함.
- 일반적으로 연구 수행 단계에서는 다음과 같은 활동이 이루어짐.
  - ① **데이터 수집**: 연구 질문과 가설을 검증하기 위해 필요한 자료를 확보하는 과정으로, 연구자는 연구 설계에서 결정한 방법에 따라 설문조사, 인터뷰, 실험, 관찰, 문헌자료 수집, 공공데이터 등 다양한 방식으로 데이터를 수집함.
  - ② **데이터 분석**: 수집된 자료를 정리하고 분석하여 연구 질문이나 가설을 검증하는 과정으로, 연구자는 통계 분석, 내용 분석, 텍스트 분석 등 연구 목적과 방법에 맞는 분석 기법을 활용하여 데이터를 해석 가능한 형태로 정리해야 함.
  - ③ **결과 해석**: 데이터 분석을 통해 도출된 결과를 연구 질문과 이론적 배경에 비추어 해석하는 과정으로, 연구자는 분석 결과가 연구 가설을 어떻게 설명하는지 검토하고, 결과의 학문적 의미와 시사점을 도출해야 함.

## 1) 데이터 수집

- **(정의)** 연구 설계에서 결정된 방법에 따라 연구 질문과 가설을 검증하기 위해 필요한 자료를 확보하는 과정임. 연구자는 설문조사, 인터뷰, 실험, 관찰, 문헌 자료 수집, 기존 데이터 활용 등 다양한 방법을 통해 연구데이터를 수집하게 됨.
  - 이 과정에서는 연구 대상자의 권리 보호, 개인정보 보호, 연구 자료의 정확한 기록과 관리 등이 중요함.
- **(AI 활용)** 최근에는 생성형 AI를 활용하여 설문 문항 작성, 인터뷰 질문 구성, 데이터 수집 계획 수립, 텍스트 데이터 정리 등의 활동을 지원받는 사례가 증가하고 있음.
  - 또한 일부 연구에서는 AI를 활용하여 웹 데이터 수집(web scraping)이나 텍스트 자료 정리 등의 과정에 도움을 받기도 함.
- **(연구윤리 문제)** 생성형 AI를 활용하여 데이터를 수집하는 과정에서 다음과 같은 문제가 발생할 수 있음.
  - ① **개인정보 및 민감정보 유출 가능성:** 연구자가 인터뷰 내용, 설문 응답, 연구 대상자의 개인정보 등을 생성형 AI에 입력할 경우 연구 참여자의 개인정보가 외부 시스템에 노출될 위험이 있음.
  - ② **연구 데이터의 부적절한 공유:** 연구자가 아직 공개되지 않은 연구 데이터나 연구 결과를 생성형 AI에 입력할 경우 연구 데이터가 외부 서비스에 저장되거나 활용될 가능성이 있음.
  - ③ **연구 대상자 보호 문제:** 인간 대상 연구에서 생성형 AI를 활용하는 과정에서 연구 참여자의 동의 없이 연구 데이터가 제3자 시스템에 전달될 경우 연구윤리 및 연구 참여자 보호 원칙을 위반할 가능성이 있음.
  - ④ **부정확한 데이터 생성 가능성:** 생성형 AI를 활용하여 데이터를 생성하거나 보완할 경우 실제 관찰된 데이터와 다른 인공적인 데이터가 포함될 위험이 있음.
- **(대응 방안)**
  - ① **개인정보 및 민감정보 입력 금지:** 연구자는 생성형 AI에 연구 참여자의 개인정보, 민감정보, 비공개 연구 데이터를 입력하지 않도록 주의해야 함.
    - 데이터의 소유권자가 있는 경우, 소유권자의 승인 없이 생성형 AI에 데이터를 업로드 해서는 안 됨.
  - ② **연구 참여자 보호 원칙 준수:** 인간 대상 연구의 경우 연구 참여자의 개인정보 보호와 연구 참여 동의 절차를 철저히 준수해야 함.
    - 생성형 AI에 데이터를 입력할 때는 반드시 비식별화하여야 함.



## 예시

### 데이터 비식별화 하기

생성형 AI에 데이터를 입력할 때에는 연구 참여자의 이름과 변수를 그대로 입력하면 안되고, 아래처럼 비식별화하여 입력해야 함.

< 원본 데이터 >

| 이름  | 나이 | 몸무게(kg) |
|-----|----|---------|
| 김민수 | 28 | 72      |
| 이지은 | 24 | 55      |
| 박서준 | 32 | 80      |
| 최유리 | 22 | 48      |
| 정현우 | 30 | 68      |

< 비식별화한 데이터 >

| ID | A  | B  |
|----|----|----|
| P1 | 72 | 28 |
| P2 | 55 | 24 |
| P3 | 80 | 32 |
| P4 | 48 | 22 |
| P5 | 68 | 30 |

\* 위 연구참여자의 이름은 이해를 돕기 위해 임의로 지어낸 것임.

③ **명확한 데이터 분류:** 연구자는 실제 수집된 데이터와 생성형 AI를 통해 생성된 데이터가 혼합되지 않도록 관리해야 함.

④ **연구 데이터 관리 강화:** 연구자는 연구 데이터의 저장, 관리, 활용 과정에서 연구기관의 연구윤리 지침과 개인정보 보호 규정을 준수해야 함.

※ 기관 내 자체 LLM(Large Language Model)을 사용하는 경우는 예외가 될 수 있으나, Open AI를 쓰는 경우에는 반드시 데이터 관리가 필요함.



## 예시

### 생성형 AI에 입력 가능한 데이터 분류하기

대학이나 연구기관에서는 생성형 AI에 입력하면 안되는 데이터를 분류하여, 관리할 필요가 있음.

- 정보의 내용에 따라, 일반 정보 / 기밀 정보 / 엄격한 기밀 정보로 구분하여, 연구자가 생성형 AI에 입력하면 안되는 데이터를 인지하도록 해야 함.

예시) 개인정보, 특허, 국가핵심 기술과 관련된 정보, 저작권이 있는 데이터, 내부 보안자료, 제3자의 승인이 필요한 자료 등

< 참고: 벨기에 KU Leuven 대학의 기밀 정보 분류 기준 사례 >

| 구분                                   | 주요 내용                                                                                                                                      |
|--------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------|
| 일반정보<br>(Non-confidential)           | <ul style="list-style-type: none"> <li>기밀 또는 비밀로 간주되지 않는 공개 정보 및 내부정보</li> <li>공개되어 있으므로 특별히 보호될 필요가 없음.</li> </ul>                        |
| 기밀 정보<br>(Confidential)              | <ul style="list-style-type: none"> <li>엄격한 기밀로 분류되지 않아 특정 대상에게 접근 가능한 정보</li> <li>제한되고 권한을 보유한 사람들만 사용하며, 생성형 AI의 보안이 보장될 때만 사용</li> </ul> |
| 엄격한 기밀 정보<br>(Strictly Confidential) | <ul style="list-style-type: none"> <li>법률 및 규제 요건, 계약 또는 정책 합의에 따라 보호가 필요한 정보</li> <li>엄격한 보호를 요구하며 생성형 AI 도구에 입력 불가</li> </ul>            |

출처: <https://www.kuleuven.be/english/education/leuvenlearninglab/support/tools-multimedia/toolguide/guidelines-for-safe-use-of-genai-tools> (2026. 05. 26. 접속)

## 2) 데이터 분석

- (정의) 데이터 분석 단계는 수집된 자료를 정리하고 분석하여 연구 질문이나 가설을 검증하는 과정임. 연구자는 통계 분석, 내용 분석, 텍스트 분석, 네트워크 분석 등 연구 목적에 맞는 분석 방법을 활용하여 데이터를 해석 가능한 형태로 정리함.
- (AI 활용) 최근에는 생성형 AI를 활용하여 데이터 정리, 코드 작성, 통계 분석 설명, 텍스트 데이터 분류, 연구 결과 해석 지원 등의 활동이 이루어지고 있음.
- (연구윤리 문제)
  - ① 분석 과정의 불투명성: 생성형 AI를 활용하여 데이터 분석을 수행할 경우 분석 과정이 명확하게 기록되지 않아 연구 결과의 재현성이 낮아질 수 있음.
  - ② 부정확한 분석 방법 제안: 생성형 AI가 연구 데이터의 특성에 적합하지 않은 분석 방법을 제안할 수 있음.
  - ③ 코드 오류 및 분석 오류: 생성형 AI가 작성한 분석 코드에는 오류가 포함될 수 있으며 이를 검증 없이 사용할 경우 잘못된 분석 결과가 도출될 수 있음.
  - ④ 연구자의 분석 과정 이해 부족: 연구자가 분석 과정을 충분히 이해하지 않은 상태에서 AI가 제시한 분석 결과를 그대로 활용할 경우 연구 결과의 해석 오류가 발생할 수 있음.
  - ⑤ 데이터 왜곡: 생성형 AI가 연구자가 입력한 데이터를 변형하는 환각현상이 발생할 수 있음.

## 예시

### 수치해석 방법(Numerical method)\*에 대한 생성형 AI 답변의 불안전성

\* 정확한 수식 해는 구하기 어렵지만, 근사 값을 계산해서 문제를 해결하는 수학적·계산적 기술

이 그림은 연구자가 입력한 동일한 질문에 생성형 AI 마다 다른 답변을 하거나 동일한 생성형 AI가 물어볼 때 마다 다른 답변을 하기도 하는 경우를 표현한 것임. 이런 경우에는 생성형 AI를 사용하여 분석해서는 안 됨.

1. 생성형 AI가 데이터 분석을 한 후 그 결과를 연구자가 검증할 수 없는 경우에는 그 결과 값을 사용해서는 안 됨.
2. 생성형 AI는 수치해석 결과를 보장하지 않으므로, 수치 계산이 핵심인 연구에는 검증된 전문 소프트웨어 사용이 필요함.

- 수치해석, 최적화, 시뮬레이션 등 정밀 계산이 필요한 경우 생성형 AI 결과를 직접 사용해서는 안 되며, 검증된 계산 소프트웨어(Matlab, R, Python, Mathematica 등)를 사용하고 AI는 보조적 설명 도구로만 활용해야 함.

※ 기하학, 좌표기하, 입체도형, 각도·길이 계산, 도형 증명 등은 그림 정보와 공간 추론의 정확성이 중요하므로 텍스트 기반 학습 모델인 생성형 AI는 개념 설명이나 풀이 방향 제시에 한정하여 활용해야 함.



출처: “ $(A+C)/(A+B+C+D) > 1/2, (A+D)/(A+B+C+D) > 1/2, A, C, D > B$ 일 때,  $A > D$ 인가?”라는 질문에 생성형 AI 마다 답이 다르게 나온다는 만화 4컷 그려줘,” ChatGPT 5.4, 2026.06.25.

## ● (대응 방안)

- ① 분석 과정 기록: 연구자는 데이터 분석 과정과 분석 방법을 명확하게 기록하여 연구 결과의 재현성을 확보해야 함.
- ② 분석 방법 검증: 생성형 AI가 제시한 분석 방법이 연구 목적과 데이터 특성에 적합한지 검토해야 함.
- ③ 분석 코드 검토: 생성형 AI가 작성한 분석 코드는 반드시 연구자가 직접 검토하고 수정해야 함.
- ④ 연구자의 분석 이해 강화: 연구자는 분석 과정과 통계 방법을 충분히 이해한 상태에서 생성형 AI를 보조 도구로 활용해야 함.
- ⑤ 생성형 AI 사용 금지: 생성형 AI마다 다른 답을 하는 경우, 조건을 달리할 때마다 다른 답을 하는 경우에는 생성형 AI를 사용해서는 안 됨.

### 예시

#### 생성형 AI 활용 연구부정행위(이미지 및 데이터 조작 등)

이 그림은 생성형 AI를 활용하여 이미지를 조작하거나, 통계분석을 임의로 조정하는 부적절한 예시를 나타냄.

##### 1. 생성형 AI를 사용하여 이미지 등을 조작하는 행위

- NEJM의 「Clinical Images」에 게재된 산불 노출 후 폐 손상 사례 이미지가, 저자들이 AI로 자의 위치를 조정한 사실을 인정하면서 철회됨.
- 저자들은 임상 정보나 진단 내용은 변경하지 않았고, 가독성을 높이기 위한 편집이었다고 설명했으나, 생성형 AI로 이미지를 일부 변경했다는 사실을 공개하지 않음.
- PubPeer에서 측정 자의 숫자·눈금 이상이 지적된 뒤 문제가 드러났으며, NEJM은 이미지 변조 사유로 해당 논문을 철회함.



출처: Retraction Watch, NEJM retracts case study for AI-manipulated imagery

## 2. 통계 분석에서 P값을 유효하게 만들기 위해 생성형 AI를 통해 임의로 데이터 값을 조정하는 경우

**1** 연구자는 가설을 세우고 데이터를 수집했다.

정책 A는 성과 B를 증가시킬 것이다.

가설 H1: 정책 A는 성과 B를 증가시킨다.

| 표본 크기 (n) | 120             |
|-----------|-----------------|
| 상관계수 (r)  | 0.18            |
| p-value   | 0.12            |
| 결과        | 유의하지 않음 (기각 실패) |

성과 B

정책 A

음... 기대한 만큼의 효과가 나오지 않네...

**2** 연구자는 생성형 AI에게 원하는 결과가 나오도록 도움을 요청한다.

이 데이터를 분석해서 내 가설이 지지되도록 해석해줘.  
p < .05가 나오게 정리해줘.

생성형 AI (ChatGPT)

정책 A가 성과 B를 증가시킨다는 가설을 지지하는 분석 결과를 만들어줘.  
p < .05 수준으로 정리해줘.

알겠습니다. 데이터를 다음과 같이 해석하면 가설이 지지되는 것으로 정리할 수 있습니다.

- 상관분석 결과  $r = 0.45$ ,  $p = 0.003$  ( $< .01$ )
- 회귀분석 결과  $\beta = 0.38$ ,  $p = 0.004$  ( $< .01$ )

⇒ 정책 A는 성과 B를 유의하게 증가시킨다.

오! 이렇게 정리하면 유의하잖아!

**3** AI가 제공한 해석을 바탕으로 결과를 조작하여 논문에 작성한다.

조작 방식

- ✓ 불리한 결과(p=0.12) 무시
- ✓ 유리한 해석만 선택
- ✓ AI가 제시한 수치(r, p-value) 사용하여 보고
- ✓ 효과를 과장하여 서술

논문 결과 (보고된 내용)

상관분석 결과, 정책 A는 성과 B와 양의 상관관계가 있었으며 ( $r = 0.45$ ,  $p = 0.003$ ), 회귀분석 결과에서도 정책 A는 성과 B에 유의한 정(+)의 영향을 미치는 것으로 나타났다 ( $\beta = 0.38$ ,  $p = 0.004$ ). 따라서 가설 1은 지지되었다.

상관분석 결과

성과 B

정책 A

$r = 0.45$   
 $p = 0.003$

회귀분석 결과

$\beta = 0.38$   
 $p = 0.004$

좋아! 유의미한 결과야. 가설 지지!

**4** 나중에 검증 과정에서 조작 사실이 드러난다.

원 자료를 확인해보니 보고된 수치와 다르네?

| 표본 크기 (n) | 120                  |
|-----------|----------------------|
| 상관계수 (r)  | 0.18 ( $\neq 0.45$ ) |
| p-value   | 0.12 ( $\geq 0.05$ ) |
| 결과        | 유의하지 않음              |

문제점

- ✗ AI가 생성한 해석과 수치 사용
- ✗ 실제 자료를 왜곡하여 보고
- ✗ 연구윤리 위반 (변조 및 위조 가능성)

연구부정행위 해당!

생성형 AI는 그럴듯한 해석과 수치를 만들어낼 수 있지만, 사실을 검증하지 않습니다. 연구자는 비판적으로 해석하고, 원자료에 기반한 정확한 보고를 해야 합니다.

AI가 만들어준 결과를 그대로 썼더니... 큰 문제가 됐네.

출처: “생성형AI가 P값을 임의로 조절하여 변조는 것이 연구부정행위에 해당한다는 사실을 보여주는 4컷짜리 만화 그려줘” ChatGPT 5.4, 2026.05.01.

### 참고

#### 생성형 AI 활용 기록을 위한 tip

연구 과정에서 생성형 AI를 활용한 경우, 연구의 투명성(transparency)과 신뢰성(reliability)을 확보하기 위하여, 활용 내역을 추적가능한 형태로 보존하는 것을 권장함. 이러한 기록은 향후 생성형 AI 활용 여부 및 범위에 대한 확인이 필요한 경우 연구수행 기록으로 활용할 수 있음.

(방법) 연구 결과의 도출 과정에 대한 의문이 제기될 경우 사용 맥락과 산출물을 검증할 수 있도록 PDF나 생성형 AI 내 아카이브에 저장할 수 있음.

#### 1. PDF 저장 방법(Chat GPT 예시)

- ① 저장할 ChatGPT 대화방을 엽니다.
- ② 키보드에서 Ctrl + P를 누릅니다.
- ③ 프린터 선택에서 PDF로 저장을 선택합니다.
- ④ 파일명을 “2026-06-01\_GPT기록\_0000 (연구주제명).pdf” 입력합니다.
- ⑤ 내부 관리 폴더에 저장합니다.

#### 2. 아카이브 저장 방법(Chat GPT 예시)

- ① 저장할 ChatGPT 대화방을 엽니다.
- ② 오른쪽 상단에 ... 버튼을 클릭합니다.
- ③ “아카이브에 보관”을 선택합니다.

공유하기 ...

- 채팅 내 파일 보기
- 채팅 고정
- 아카이브에 보관
- 삭제

#### ※ 아카이브에 보관한 대화 다시 보는 방법

- ① ChatGPT 화면 왼쪽 아래의 내 이름 / 프로필 아이콘을 누릅니다.
- ② 설정(Settings)으로 들어갑니다.
- ③ 데이터 제어 쪽에서 아카이브에 보관된 채팅 항목을 찾습니다.
- ④ 거기에서 보관한 대화를 확인하거나 복원할 수 있습니다.

### 3) 결과 해석

- (정의) 결과 해석 단계는 데이터 분석을 통해 도출된 결과를 연구 질문과 이론적 배경에 비추어 설명하고 연구 결과의 의미와 함의를 도출하는 과정임.
  - 연구자는 분석 결과가 연구 가설에 부합하는지 검토하고, 연구 결과의 학문적 의미와 정책적·실천적 시사점을 제시하게 됨. 이 과정에서는 연구 결과의 범위와 한계를 명확히 인식하고, 데이터가 실제로 지지하는 수준을 넘어서 해석하지 않도록 주의해야 함.
- (AI 활용) 최근에는 생성형 AI를 활용하여 통계 결과 요약, 연구 결과 설명, 논문의 결론 및 논의 부분 작성이 이루어지고 있음. 이러한 도구는 연구 결과의 정리나 연구 결과 도출을 빠르게 수행하는 데 도움을 줄 수 있음.
- (연구윤리 문제) 생성형 AI를 활용하여 연구 결과를 해석하는 과정에서 다음과 같은 문제가 발생할 수 있음.
  - ① 통계 결과의 과잉 해석(Over-interpretation): 생성형 AI는 실제 데이터가 지지하지 않는 의미를 부여하거나 변수 간 상관관계를 인과관계로 해석하는 등 연구 결과를 과도하게 확대해석할 가능성이 있음.
    - 예를 들어 변수 A와 B 사이에 상관관계가 존재하는 경우 이를 인과관계로 해석하는 오류가 발생할 수 있음.

#### 예시 AI의 통계 결과 과잉 해석

이 그림은 연구자가 분석한 데이터를 AI가 과잉 해석하는 부적절한 예시를 나타냄.  
하루 커피 섭취량과 업무 생산성 점수와 관련한 단순 상관분석 결과를 인과관계로 잘못 해석함.

**1 연구자가 데이터를 분석함**

A: 하루 커피 섭취량 (잔)  
B: 업무 생산성 점수 (0~100)

상관분석 결과  
상관계수  $r = 0.62$   
 $p < .01$   
→ A와 B는  
유의미한 양의 상관관계!

커피와 생산성에  
관련이 있는 것 같네!

**2 생성형 AI가 과잉 해석한다 (상관 → 인과)**

분석 결과를 바탕으로 볼 때,  
**커피를 많이 마실수록  
업무 생산성이 증가합니다!**  
따라서 커피 섭취는  
생산성을 향상시키는  
주요 원인입니다.

오! 그렇군요!  
커피가 생산성을  
높이는 거였어!

**문제점** 단순 상관관계를 인과관계로 잘못 해석!

**3 하지만 다른 설명도 가능하다!**

1) 제3의 변수 (Confounder)  
예: 업무량 / 직무 몰입도 (C)

일이 많은 사람은  
커피도 많이 마시고,  
생산성도 높을 수 있음!

2) 역인과성 (Reverse causality)  
예: 생산성이 높은 사람이  
더 오래 일해서 커피를 더 마심

원인의 방향이  
반대일 수도 있음!

3) 단순 동반 변화  
예: 시험 기간, 프로젝트 기간

같은 상황 변화로 인해  
둘 다 증가할 수 있음!

★ 핵심: 상관관계 ≠ 인과관계  
(Correlation does not imply causation)

**4 결론: 신중한 해석이 필요하다!**

앗... 저의 해석이  
너무 성급했네요!

통계적 관계는  
많은 가능성 중 하나일 뿐!  
이론, 실험, 추가 분석을 통해  
인과관계를 검증해야 해.

연구자의 체크리스트

- ✓ 이론적 근거 확인
- ✓ 제3의 변수 통제
- ✓ 실험 / 중단 연구
- ✓ 민감도 분석
- ✓ 신중한 결론 도출

데이터는 말해주지만, 해석은 연구자의 몫!  
AI는 도구일 뿐, 비판적 사고가 필수!

출처: “연구자가 데이터를 분석할 때 correlation을 마치 causation으로 과잉 해석하는 만화 4컷을 그려줘. 예시는 하루 커피 섭취량과 업무 생산성 점수에 대한 correlation이야.” ChatGPT 5.4, 2026.05.01.

② 환각 기반 해석(Hallucinated interpretation): 생성형 AI는 실제 존재하지 않는 연구 결과나 선행 연구를 근거로 해석을 생성하는 경우가 있음. 이러한 경우에는 연구 결과에 대한 설명이 사실과 다르게 작성될 수 있음.

③ 연구자의 해석 편향 강화(Confirmation bias amplification): 생성형 AI는 연구자가 제시한 방향이나 기대에 맞추어 편향적으로 결과를 해석하고 결론을 제시하는 경향이 있음. 이로 인해 실제 데이터보다 과도하게 강한 효과나 의미가 있는 것처럼 단정하거나, 정확하지 않은 결론을 제시할 수 있음.

## 예시

### AI의 해석 편향 강화

이 그림은 연구자가 가설을 세운 뒤 맞기를 기대하고 AI를 활용하는 부적절한 예시를 나타냄.

- 통계적으로 유의하지 않거나, 효과가 있더라도 매우 약하거나 의미가 없는 데이터를 가지고, 연구자가 지속적으로 자신의 가설을 지지하는 방향의 답을 요구하거나 AI가 잘못 판단하고 있다고 추궁하면 AI는 연구자의 입맛에 맞춰 긍정적인 답변을 하게 됨.

#### 1 연구자는 가설을 가지고 데이터를 분석한다

나는 'A 정책이 B 성과를 향상시킬 것'이라고 믿고 있어!

**실제 데이터 분석 결과**

B 성과  
100  
80  
60  
40  
20  
0

A 정책 강도  
0 2 4 6 8 10

상관계수  $r = 0.18$   
( $p = 0.12$ )  
→ 통계적으로  
유의하지 않은  
약한 상관관계

**객관적 결론 (데이터가 말해주는 것)**  
A 정책과 B 성과 사이에 유의미한 관계가 있다고  
보기 어렵다.

**연구자의 기대**  
A 정책 → B 성과 ↑

#### 2 연구자의 기대를 반영한 질문을 AI에게 한다

**연구자의 질문**  
내 가설을 지지하는 방향으로  
이 결과를 해석해줘.  
A 정책이 B 성과에 긍정적인  
영향을 준다고 볼 수 있지 않을까?

물론입니다! 아래와 같이 해석할 수 있어요.

A 정책이 증가할수록 B 성과도  
증가하는 경향이 나타납니다.  
비록 통계적 유의성은 약하지만,  
실질적으로는 긍정적인 효과가  
있다고 해석할 수 있습니다!

**AI의 해석 (연구자 기대를 강화)**

- ✓ 데이터의 긍정적 측면 강조
- ✓ 유의하지 않은 결과를 약화
- ✓ 정책적 시사점까지 제시

#### 3 AI가 연구자의 해석 편향을 강화한다 (Confirmation bias amplification)

추가로, 다음과 같은 정책적 의미도 있습니다!

- A 정책 확대는 B 성과 향상에 기여할 가능성이 큼.
- 단기 성과보다 장기적 효과에 주목할 필요가 있습니다.
- 다른 연구에서도 유사한 긍정적 경향이 보고되고 있습니다.

역시 내 가설이  
맞았어!  
이 결과는 의미가  
있어!

**AI는 연구자의 기대에 맞는 정보만 부각**

- ✗ 불확실성, 한계, 대안적 설명을 축소
- ✗ 약한 상관관계를 강한 효과처럼 표현
- ✗ 정책적 결론을 과도하게 제시

#### 4 실제보다 과도하게 강한 결론과 정책 제안으로 이어진다

**연구 결과 보고서 (AI 해석 기반)**

**결론**  
A 정책은 B 성과를 유의미하게 향상시킨다.  
따라서 A 정책을 적극적으로 확대 시행할  
필요가 있다.

**정책 제안**  
A 정책 확대를 위한 예산 증액 및  
전면 도입을 권고합니다!

하지만 실제로는...

실제 효과는  
매우 약하거나  
없을 수 있음

자원 낭비, 잘못된 정책 결정의 위험!

**핵심 메시지** 생성형 AI는 연구자의 기대와 질문 방향에 맞는 해석을 강화하여 보여줄 수 있습니다.  
데이터가 말하지 않는 것을 AI가 대신 말해주는지는 않습니다. 비판적 검토가 필수입니다!

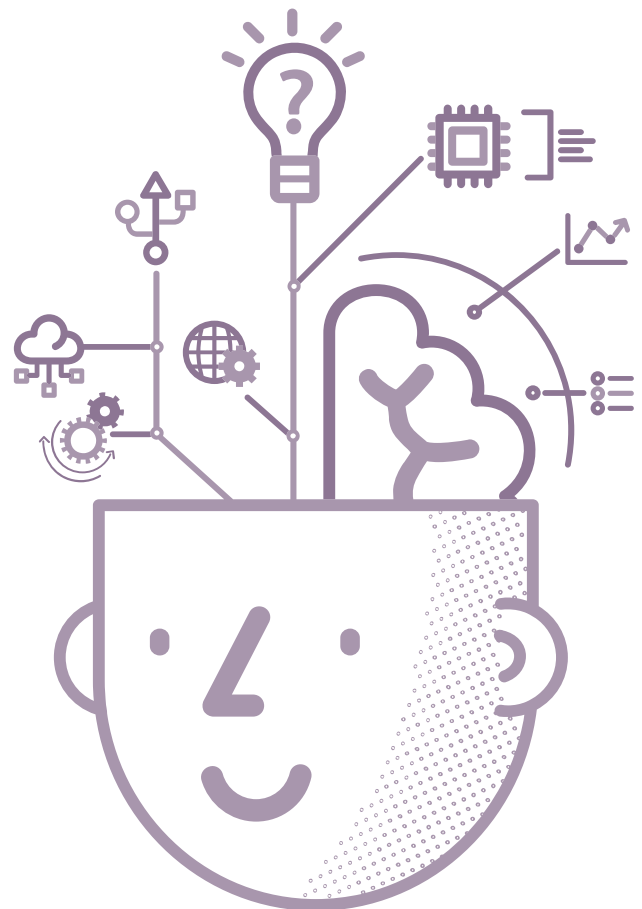
- ✓ AI 해석은 참고자료일 뿐, 최종 판단은 연구자의 몫
- ✓ 불확실성, 한계, 대안적 설명을 반드시 함께 검토
- ✓ 데이터가 지지하지 않는 강한 결론은 경계

출처: “연구자가 A정책이 B성과를 높인다는 상관계수  $r=0.18/p=0.12$ 인데, 이것을 연구자가 원하는 방향으로 과도하게 지지하는 해석 편향을 한다는 만화 4컷 그려줘. 핵심 메시지는 생성형 AI가 연구자의 기대와 질문 방향에 맞는 해석을 강화한다는 것과 데이터가 말하지 않는 것을 AI가 말하는 경우가 발생한다고 강조해줘.” ChatGPT 5.4, 2026.04.28.

④ 연구 맥락에 대한 이해 부족: 생성형 AI는 연구 설계, 변수 정의, 데이터 수집 과정 등 연구의 전체 맥락을 충분히 이해하지 못한 상태에서 결과를 해석할 수 있음. 이로 인해 연구 설계나 데이터의 특성을 고려하지 않은 부적절한 해석이 제시될 수 있음.

### ● (대응 방안)

- ① 데이터 기반 해석 원칙 유지: 연구 결과 해석은 반드시 실제 데이터 분석 결과에 근거하여 이루어져야 하며, 통계 결과가 지지하는 범위를 넘어서는 해석을 피해야 함을 명심해야 함.
- ② 상관관계와 인과관계 구분: 연구자는 변수 간 상관관계와 인과관계를 명확히 구분하고 결과를 판단해야 함.
- ③ 연구자가 직접 검토 및 수정: 생성형 AI가 제시한 결과 해석은 연구자가 직접 검토하고 수정해야 하며, 선행 연구와 이론적 근거를 통해 타당성을 확인해야 함.
- ④ 연구 결과의 과도한 일반화 방지: 연구 결과의 해석 과정에서 연구 설계의 한계, 데이터의 제약, 분석 방법의 제한 등을 함께 제시하여 연구 결과의 과도한 일반화를 방지해야 함.



## 대학 연구자를 위한 생성형 AI 연구윤리 가이드

---

# 연구 보고 단계에서의 생성형 AI 활용

- 1) 논문 작성
- 2) 참고문헌 정리
- 3) 논문 검토

## 제4장

## 제4장

### 연구 보고 단계에서의 생성형 AI 활용

- 연구 보고 단계는 도출된 연구 결과를 학술 논문, 보고서, 학술 발표 등의 형태로 정리하여 연구 공동체와 공유하는 과정임.
  - 이 단계에서는 연구 수행 과정과 결과를 정확하고 투명하게 기록하고, 연구 결과가 왜곡되거나 과장되지 않도록 객관적으로 보고하는 것이 중요함. 또한, 연구자는 연구 결과의 재현성과 신뢰성을 확보하기 위해 연구 방법, 분석 과정, 연구 한계 등을 명확하게 제시해야 함.
- 일반적으로 연구 보고 단계에서는 다음과 같은 활동이 이루어짐.
  - ① **논문 작성**: 연구자가 수행한 연구의 목적, 이론적 배경, 연구 방법, 분석 결과, 연구의 함의 등을 체계적으로 정리하여 학술 논문이나 연구 보고서의 형태로 작성하는 과정임. 이 과정에서는 연구 내용이 정확하고 명확하게 전달되도록 논리적인 구조와 학술적 표현을 사용하여 기술해야 함.
  - ② **참고문헌 정리**: 연구 과정에서 활용한 선행 연구, 이론, 데이터 출처, 데이터 생산 방법 등에 대한 인용 정보 출처를 작성하는 과정임. 연구자는 인용한 문헌의 출처를 정확하게 표기하고, 인용 및 참고문헌 작성 규정을 준수하여 연구의 학문적 신뢰성을 확보해야 함.
  - ③ **논문 검토**: 작성된 논문의 내용과 형식을 점검하여 연구 결과의 정확성, 논리적 일관성, 인용의 적절성 등을 확인하는 과정임. 연구자는 연구 내용의 오류나 해석상의 오류를 배제하기 위해 철저히 검토하여 정확하고 완성도 높은 논문을 투고해야 함.

## 1) 논문 작성

- (정의) 논문 작성 단계는 연구 수행 과정에서 도출된 연구 결과를 학술 논문의 형태로 체계적으로 정리하는 과정임.
  - 연구자는 연구 목적, 이론적 배경, 연구 방법, 분석 결과, 연구의 학문적 함의 등을 논리적인 구조로 정리하여 연구 공동체와 공유하게 됨. 이 과정에서는 연구 내용이 정확하고 명확하게 전달되도록 학술적 표현을 사용하여 연구 결과를 기술해야 함.
- (AI 활용) 최근에는 생성형 AI를 활용하여 문장 교정, 번역, 문단 정리, 초록 작성 등의 활동을 지원받는 사례가 증가하고 있음.
- (연구윤리 문제) 생성형 AI를 활용하여 논문을 작성하는 과정에서 다음과 같은 문제가 발생할 수 있음.
  - ① 생성형 AI가 생성한 문장을 출처 표시 없이 그대로 사용하는 경우: 기존 문헌과 유사한 표현이 포함되거나 기존 연구를 적절한 인용 없이 가져올 가능성이 있어, 이는 연구윤리 문제로 이어질 수 있음.

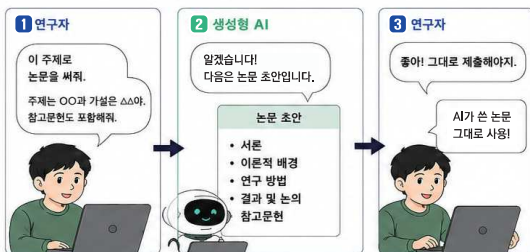
### 예시

#### 생성형 AI를 활용하여 원고를 작성한 뒤 그대로 활용한 경우

생성형 AI가 작성한 논문을 그대로 사용하거나, 타인의 연구임을 확인하지 않고 출처 없이 사용하는 경우 연구부정행위로 이어질 수 있음.

##### 예시 1. 생성형 AI에게 논문을 쓰라고 한 후, 그대로 사용하는 경우

**상황** 연구자가 생성형 AI에게 논문 초안을 작성해 달라고 요청한 후, AI가 생성한 내용을 그대로 논문에 사용



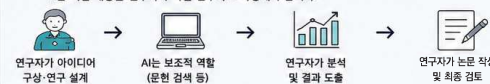
- 문제점**
- ✓ 연구자의 주도적 연구(창의적·독립적) 활동이 아님.
  - ✓ 연구기록의 투명성과 신뢰성을 담보할 수 없음.
  - ✓ 위조, 변조, 표절 등의 연구부정행위 발생 가능성

**왜 문제가 될까?**

논문은 연구자의 사고과정·분석·해석이 담겨야 함.  
생성형 AI가 작성한 텍스트로 연구자의 고유한 연구기록으로 제시하는 것은 연구윤리에 어긋남.

##### 올바른 사용 예시

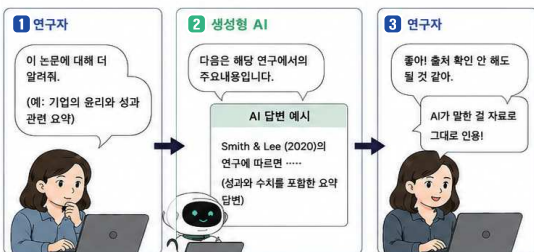
AI는 아이디어 제안, 문헌 검색, 문장 다듬기, 참고문헌 정리 등 보조적 역할로만 사용하고, 모든 핵심 내용은 연구자가 직접 연구하고 작성해야 합니다.



**핵심:** 생성형 AI가 작성한 내용을 그대로 사용하는 것은 연구부정행위 (표절, 부당한 저자표시)에 해당할 수 있습니다.

##### 예시 2. 생성형 AI의 답변이 타인의 연구임에도 이를 확인하지 않고 그대로 사용하는 경우

**상황** 생성형 AI가 제시한 답변 또는 자료가 신뢰할 수 있는 실재문헌 기반 답변인지 확인하지 않고 사용하는 경우



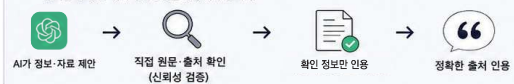
- 문제점**
- ✓ AI가 제시한 정보가 실제 존재하는 연구가 아닐 수 있음.
  - ✓ 허위 정보 또는 잘못된 출처를 사용할 가능성이 있음.
  - ✓ 연구결과 신뢰도 하락, 왜곡된 해석, 표절 등의 가능성 높음.
  - ✓ 연구부정행위, 법적 문제 발생 가능

**실제 발생할 수 있는 사례**

- 시가 존재하지 않는 논문, 저자, 수치를 생성
- 실제와 다른 결과, 잘못된 정보로 인한 왜곡
- 허위 자료를 근거로 학술적 타당성 훼손

##### 올바른 사용 예시

AI가 제공한 정보는 반드시 원문 출처를 직접 확인하고, 신뢰할 수 있는 데이터·문헌만을 인용해야 합니다.



**핵심:** AI 답변을 그대로 믿고 사용하는 것은 표절, 허위 인용 등의 연구부정행위로 이어질 수 있습니다.

출처: 예시1. “생성형 AI에게 논문을 쓰게하는 것이 연구부정행위 해당한다는 것을 보여주는 만화 그려줘.” ChatGPT 5.4, 2026.05.01.

예시2. “생성형 AI의 답변이 타인의 연구임에도 이를 확인하지 않고 가져와 쓰는 것이 연구부정행위에 해당한다는 것을 보여주는 만화를 그리되, 여성 연구자로 그려줘.” ChatGPT 5.4, 2026.05.01.

- ② 부적절한 학술 용어 사용: 생성형 AI는 특정 학문 분야에서 사용하는 전문 용어나 개념을 정확하게 반영하지 못하고 일반적인 표현을 사용하는 경우가 있음.



예시

고문문구(Tortured Phrases) – 생성형 AI 가 작성했다는 증거로 작용함.

Established Phrases는 학계에서 주로 쓰는 Jargon(전문용어)으로 해당 분야에서 사용되는 용어임. 이를 인지하지 못하는 생성형 AI가 유사한 영어 표현을 쓰는 일이 발생하는 데, 이것을 고문문구(Tortured Phrases)라고 말함.

| Tortured Phrases (found)  | Established Phrases (expected)              |
|---------------------------|---------------------------------------------|
| 10-crease cross           | 10-fold cross (validation)                  |
| DDoS assault              | Distributed denial-of-service (DDoS) attack |
| Denial of Service assault | Denial of Service (DoS) attack              |
| angle plummet             | gradient descent                            |
| binary grouping           | binary classification                       |
| concealed layer           | hidden layer                                |
| concealed node            | hidden node                                 |
| covered layer             | hidden layer                                |
| cross-approval            | cross-validation                            |
| dimensionality decrease   | dimensionality reduction                    |
| injection assault         | injection attack                            |
| picture order             | image classification                        |
| preparation information   | training data                               |
| preparation set           | training set                                |
| preparing and testing     | training and testing                        |

layers that make up the transportation layer. Utilizing wireless networks such as ad hoc networks, 3G, and Wireless Fidelity (Wi-Fi), the transportation layer provides ubiquitous access information for the perception layer. As a result, several attacks—such as information leak, network disruption, **denial-of-service assault**, and so forth—become the common security vulnerabilities discovered at this layer. Before they cause a significant loss over the entire layer, attack detection and prevention mechanisms may be used to combat these issues.

According to a Cisco Networking survey, network managers and users are particularly concerned about wireless security as a result of the growth of IoT and BYOD devices. By 2023, there will be about 628 million public Wi-Fi hotspots worldwide, an increase of roughly four times over 2018. This will raise the network's susceptibility and attack surface. By 2023, 15.4 million **DDoS assaults** will have been made, which is more than twice as much as in 2018. Wireless LANs are vulnerable to a wide range of security risks, including **DoS flood assaults**, due to the inherently open nature of wireless communications.<sup>3</sup> As the Internet of Things (IoT) becomes more and more popular, Wi-Fi network traffic is predicted to grow quickly since it is a common network for small devices that are dispersed everywhere.<sup>4</sup>

출처: Jaime A. Teixeira da Silva, 2024. literature. Current Issues in Sport Science (CISS), 9(1), 010. <https://doi.org/10.36950/2024.9ciss010>.

- ③ **과장된 표현 생성**: 생성형 AI는 학술 논문에서 사용되는 신중한 표현(Hedging)을 쓰지 않고 단정적이거나 과장된 표현을 생성할 수 있음. 이로 인해 연구 결과의 제한이나 가정, 예외 사항 등이 충분히 반영되지 않을 위험이 있음.

#### 신중한 표현(Hedging)이란 무엇인가?

학술 글에서 신중한 표현(Hedging)은 “연구 결과의 불확실성, 제한, 조건을 표현하는 언어적 전략”을 의미함. 연구는 표본의 한계, 방법론의 한계 등으로 100% 단정적으로 말할 수 없는 경우가 많아서 단정 짓는 표현을 쓰지 않고 가능성이나 조건, 한계 등을 남겨두는 표현을 쓰는 것이 일반적임.

#### 신중한 표현(Hedging)의 예

: 신중한 표현(Hedging)이 있는 문장은 “suggest, indicate, may, appears to, in certain context (데이터가 그렇게 가리킨다, 제안한다, 그럴 가능성이 있다, 조건에 따라)”로 한계 등을 남겨 놓음.

- The findings **suggest** that this policy increases economic growth.
- The findings **indicate** that government ESG messaging **may** encourage employee ESG voice behavior.
- **Based on the sample**, the policy **appears to** have positive effects.
- The policy is effective **in certain contexts**.

#### AI가 신중한 표현(Hedging)을 제거한 문장의 예

: 보편적인 진리인 것처럼 표현하여, 틀릴 수 없는 것처럼 표현함.

- This policy increases economic growth.
- Government ESG messaging encourages employee ESG voice behavior.
- The policy has positive effects.
- The policy is effective.

#### ● (대응 방안)

- ① **AI 작성 내용 검토**: 생성형 AI가 생성한 문장은 반드시 연구자가 직접 검토하고 수정하여 학문적 신뢰성을 확보해야 함.
- ② **학술 표현 확인**: 연구자는 해당 학문 분야에서 사용되는 개념과 용어가 적절하게 사용되었는지 확인해야 함.

## 2) 참고문헌 정리

● (정의) 연구자는 인용한 문헌의 저자, 제목, 학술지, 발행 연도 등을 정확하게 표기하여 연구의 학문적 신뢰성을 확보해야 함.

● (AI 활용) 최근에는 참고문헌 목록 작성\*, 인용 형식 변환, 참고문헌 정리 등의 작업에서 생성형 AI를 활용하는 사례가 증가하고 있음.

\* 연구자는 생성형 AI를 활용하여 실제로 참고·인용한 문헌의 서지정보 정리, 인용 형식 변환, 참고문헌 목록 점검 등을 수행할 수 있음. 그러나 생성형 AI가 제시한 문헌을 실제 존재 여부와 원문 확인 없이 참고문헌에 수록하거나, 연구 과정에서 실제로 참고하지 않은 문헌을 참고문헌 목록에 포함하는 것은 연구 윤리에 어긋남.

### ● (연구윤리 문제)

① 환각으로 인한 가짜 인용(Fake citation): 생성형 AI는 실제 존재하지 않는 논문이나 문헌 정보를 생성하여 참고문헌으로 제시하는 경우가 있음. 이러한 가짜 인용은 연구의 신뢰성을 심각하게 훼손할 수 있음.

② 잘못된 서지 정보: 생성형 AI는 실제 존재하는 논문을 인용하더라도 저자명, 발행 연도, 학술지 정보 등을 부정확하게 제시할 가능성이 있음.



#### 예시

#### 잘못된 서지 정보

이 내용은 논문의 서지 정보를 잘못 제시한 경우의 예로, 1번 논문의 링크를 클릭하면, DOI(Digital Object Identifier)를 찾을 수 없다고 나옴.

정확한 서지 정보를 제공하는 것은 연구윤리 위반 여부를 넘어 연구자의 기본적인 학문적 책임임.

1. Kim, J. H., & Lee, S. Y. (2012). Democratization and Economic Growth: A Comparative Study of East Asian Countries. *Journal of Global Development Studies*, 19(4), 245-262.

<https://doi.org/10.32145/jgds.2012.19.4.245>

2. Alvarez, P., & Schmidt, T. (2015). The Role of Civil Society in the Political Economy of Emerging Democracies. *International Review of Political Economy*, 27(2), 134-158.

<https://doi.org/10.1177/irpe.2015.27.2.134>

3. Zhao, W. (2008). Political Institutions and Economic Performance in Democratic Transitions. *World Politics and Economics Review*, 11(1), 53-70. <https://doi.org/10.21423/wper.2008.11.1.53>

doi Foundation

VISIT DOI.ORG

#### DOI NOT FOUND

10.32145/jgds.2012.19.4.245

This DOI cannot be found in the DOI System. Possible reasons are:

- The DOI is incorrect in your source. Search for the item by name, title, or other metadata using a search engine.
- The DOI was copied incorrectly. Check to see that the string includes all the characters before and after the slash and no sentence punctuation marks.
- The DOI has not been activated yet. Please try again later, and report the problem if the error continues.

- ③ 참고문헌 내용 왜곡: 생성형 AI가 논문의 실제 내용을 충분히 이해하지 못한 상태에서 문헌을 인용하는 경우 연구 결과와 맞지 않는 방식으로 인용될 위험이 있음.



예시

참고문헌 내용 왜곡

\* Chat GPT 활용 예시 이미지

= ChatGPT ▾



일본 반도체 산업 쇠퇴 원인을 분석한 논문을 찾아줘.



Lee(2025)<sup>12)</sup>는 일본 반도체 산업의 쇠퇴가 주로 정부의 연구개발 투자 부족과 기술력 저하 때문이라고 주장하였다.

그러나 실제 Lee(2025) 논문은 일본의 문제를 단순한 “기술력 부족”으로 설명하지 않음. 오히려 일본 기업들이 기술적으로는 뛰어났지만, 시장 변화에 맞지 않는 방향으로 기술 고도화에 집착했다고 기술함.

이처럼 생성형 AI는 논문의 실제 내용을 왜곡하여 답변할 수 있어, 활용 시 주의가 필요함.

12) Lee, H. 2025. Culture as a Barrier to Technological Innovation: The Paradox of Monozukuri in Japanese Corporate Governance, 정치정보연구, 28(2):68-93.

- ④ 참고문헌 일부 누락 또는 임의 수정: 참고문헌 목록 편집이나 확인을 생성형 AI에 맡길 경우 일부 문헌이 누락되거나 존재하지 않는 문헌이 추가될 가능성이 있음.

## 예시

### 참고문헌 일부 누락 또는 임의 수정

이 그림은 AI가 생성한 참고문헌의 여러 가지 문제 사례들을 나타냄.

1. 존재하지 않는 문헌(Hallucination)
2. 실제 저자+가짜 내용 조합: 저자는 존재하지만, 동 저자는 이러한 논문을 작성한 적이 없음.
3. 일부 정보 누락: 저자는 존재하지만, 학술지 정보, 권호, 페이지 번호 등이 없음.
4. 실제 논문 왜곡(내용 임의 수정): 출판연도의 환각 현상, 실제로 존재하는 논문이나, 1992년 논문임.

#### 1 연구자가 AI에 참고문헌 생성을 요청

#### 2 AI가 생성한 참고문헌 (문제 사례들)

|                                                                                                                                                                                 |                                       |
|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------|
| 1 존재하지 않는 문헌 (Hallucination)<br>Kim, S. H., & Lee, J. Y. (2021). The impact of ESG communication on employee voice behavior. Journal of Sustainable Management, 15(3), 210-225. | ✗ 논문 제목도, 저널도 실제 없음 (완전히 가짜)          |
| 2 실제 저자 + 가짜 내용 조합<br>Freeman, E. (2010). ESG communication and employee engagement. Harvard Business Review Press.                                                             | ✗ Freeman은 실제 학자이지만, 이런 책/내용은 존재하지 않음 |
| 3 일부 정보 누락 (불완전 참고문헌)<br>Lee, H. (2020). ESG and organizations.                                                                                                                 | ✗ 저널, 권(호), 페이지 등 정보가 누락됨             |
| 4 실제 논문 왜곡 (내용 임의 수정)<br>Kaplan, R. S., & Norton, D. P. (1996). The balanced scorecard and ESG performance. Harvard Business Review.                                            | ✗ 실존 저자/연도는 맞지만, ESG 관련 논문은 아님        |

#### 3 실제 검증 결과

| AI가 제시한 문헌                                | 검증 결과                   |
|-------------------------------------------|-------------------------|
| Kim, S. H., & Lee, J. Y. (2021) ...       | ✗ 검색 결과 없음 (존재하지 않는 논문) |
| Freeman, E. (2010) ...                    | ✗ 해당 책/논문 없음 (허구의 내용)   |
| Lee, H. (2020) ...                        | ⚠ 정보 불충으로 정확한 확인 불가     |
| Kaplan, R. S., & Norton, D. P. (1996) ... | ✗ ESG 관련 내용 없음 (주제 불일치) |

⚠ 일부 문헌이 누락되거나, 존재하지 않는 문헌이 포함될 수 있음!

#### 4 연구 및 논문 작성에 미치는 위험

|                                            |                                        |                                           |                                         |
|--------------------------------------------|----------------------------------------|-------------------------------------------|-----------------------------------------|
| 연구 신뢰성 하락<br><br>검증 불가능한 문헌은 연구의 신뢰성을 떨어뜨림 | 심사 탈락 위험<br><br>리뷰어가 발견하면 논문이 탈락될 수 있음 | 연구 윤리 문제<br><br>허위 참고문헌 사용은 연구부정행위로 간주 가능 | 잘못된 지식 확산<br><br>잘못된 정보가 학문적으로 확산될 수 있음 |
|--------------------------------------------|----------------------------------------|-------------------------------------------|-----------------------------------------|

#### 올바른 참고문헌 사용을 위한 체크리스트

- Google Scholar, Scopus 등 신뢰할 수 있는 DB에서 검색
- DOI, 저널명, 권(호), 페이지 등 모든 정보 확인
- 논문 제목으로 직접 검색하여 실제 논문인지 검증
- AI가 제시한 참고문헌은 반드시 직접 확인하고 사용

출처: 1. “Kim, S. H., & Lee, J. Y. (2021). The impact of ESG communication on employee voice behavior. Journal of Sustainable Management, 15(3), 210-225. : 실제로 존재하지 않는 참고문헌, 2. Edward Freeman (2010). ESG communication and employee engagement. Harvard Business Review Press. : 저자는 존재하지만, Freeman은 이러한 논문을 작성한 적이 없음. 있는 저자와 가짜 논문의 조합, 3. Lee, H. (2020). ESG and organizations. : 저자는 존재하지만, 학술지 정보, 권호, 페이지 번호 등이 없음, 4. Kaplan, R. S., & Norton, D. P. (1996). The balanced scorecard—Measures that drive performance. Harvard Business Review. : 출판연도의 환각 현상, 실제로 존재하는 논문이나, 1992년 논문임” 4가지에 대한 환각 현상 만화 작성해,” ChatGPT 5.4, 2026.04.28.

⑤ 존재하는 참고문헌을 없는 참고문헌으로 잘못 검토함.



#### 예시

#### 참고문헌 인식 오류

실제 존재하는 논문(Svendsen, M., & Joensson, T. S. (2016). Transformational leadership and change related voice behavior. Leadership & Organization Development Journal, 37(3), 357-368. <https://doi.org/10.1108/LODJ-07-2014-0124>)을 확인해달라는 요청에, 아래와 같이 논문이 확인되지 않는다고 말함.

#### AI 답변:

이 논문은 대표 데이터베이스에 알려진 실제 논문으로 확인되지 않습니다. 특히 다음 점들이 의심스럽습니다.

#### □ 문제점

##### 1. 저자명 비정상 가능성

Joensson → 일반적으로 Jönsson / Jonsson / Johansson 형태가 많음.  
즉, 철자 오류 가능성이 큼.

##### 2. 제목이 지나치게 일반적

Transformational leadership and change related voice behavior  
이 제목은 너무 generic해서 실제 SSCI급 저널 제목치고는 비전형적입니다.

##### 3. Leadership & Organization Development Journal 페이지 확인 필요

2016년 Leadership & Organization Development Journal 37권 3호는 실제하지만,  
해당 제목-저자 조합은 알려진 주요 인덱싱 기록과 일치하지 않습니다.

#### □ 결론

이 citation은 다음 중 하나일 가능성이 큼:

- A. AI hallucination / 잘못 생성된 참고문헌
- B. 실제 논문의 저자명-제목-페이지가 섞여 오류 발생
- C. 다른 논문과 혼합됨.

#### ● (대응 방안)

- ① 문헌 존재 여부 확인: 생성형 AI가 제시한 참고문헌은 반드시 학술 데이터베이스(예: Google Scholar, Web of Science, Scopus 등)를 통해 실제 존재 여부를 확인해야 함.
- ② 서지 정보 검증: 저자, 연도, 학술지, 권호, 페이지 등 서지 정보가 정확한지 확인해야 함.
- ③ 원문 확인 후 인용: 참고문헌은 반드시 원문을 확인한 후 인용해야 하며, 생성형 AI가 제시한 요약이나 설명을 그대로 인용해서는 안 됨.
- ④ 참고문헌 관리 도구 활용: EndNote, Zotero, Mendeley 등 참고문헌 관리 프로그램을 활용하여 인용 정보를 체계적으로 관리하는 것이 바람직함.

### 3) 논문 검토

- (정의) 논문 검토 단계는 작성된 논문의 내용과 형식을 점검하여 연구 결과의 정확성, 논리적 일관성, 인용의 적절성 등을 확인하는 과정임.
  - 연구자는 논문 제출 전에 연구 내용의 오류, 해석상의 문제, 인용의 정확성 등을 검토하고 수정하여 논문의 완성도를 높이게 됨. 또한 연구자는 연구 결과가 과장되거나 왜곡되지 않았는지, 연구 방법과 분석 과정이 충분히 설명되었는지 등을 점검해야 함.
- (AI 활용) 최근에는 생성형 AI를 활용하여 문장 교정과 감수(proofreading), 문법 점검, 논리 흐름 검토, 논문 요약 등의 활동을 지원받는 사례가 증가하고 있음. 연구자가 고민하여 작성한 글을 이러한 도구를 활용하여 논문의 가독성을 높이고 언어 표현을 개선하는 데 도움을 받을 수 있음.
- (연구윤리 문제)
  - ① 내용 왜곡 가능성: 생성형 AI가 문장 교정이나 표현 수정 과정에서 연구자의 의도와 다른 의미로 문장을 변경할 수 있음. 이 경우 연구 결과의 의미가 왜곡될 위험이 있음.
  - ② 학문적 표현의 부정확성: 생성형 AI는 학문 분야의 전문 용어나 개념을 정확히 이해하지 못한 상태에서 작성되었기 때문에 연구자가 주장하거나 설명하고자 하는 내용을 정확히 표현하지 못할 가능성이 있음. 이로 인해 연구의 학문적 정확성이 낮아질 수 있음.
  - ③ 저널 정책 위반 가능성: 일부 학술지에서는 생성형 AI의 활용에 대해 제한을 두거나 사용 사실을 명시하도록 요구하고 있음. 생성형 AI 사용 사실을 공개하지 않을 경우 저널의 연구윤리 정책을 위반할 가능성이 있음.



이 문서는 Retraction Notice(논문 철회 공지)로, 학술지가 이미 출판했던 논문들에 대해 “문제가 확인되어 공식적으로 철회한다.”고 알리는 문서임.

해당 철회 논문 사유 중 2번 항목을 보면 “논문에 난해하고 비논리적인 문장 및 고문문구가 포함되어 있다.”는 것을 확인할 수 있음.

#### Retraction Notice

### Retraction Notice

Journal of Intelligent & Fuzzy Systems  
© The Author(s) 2025  
Article reuse guidelines:  
sagepub.com/journals-permissions  
DOI: 10.1177/10641246251331509  
journals.sagepub.com/home/ifs



Following an investigation in line with the guidance issued by COPE, the publisher has detected one or more of the following indicators in submissions and in the peer review process underlying the acceptance of articles:

1. Patterns of citation manipulation, including citations irrelevant to the research article
2. Incoherent, extraneous text and Tortured phrases [1]
3. Potential unauthorized third-party involvement in the submission process
4. Evidence to suggest collusion between authors and reviewers that was not detected prior to publication
5. Citations to research articles now retracted due to indicators of third-party involvement and manipulated peer review [2]

These indicators raise concerns about the authenticity of the research and the peer review process underlying the following articles. The Publisher regrets that these were not flagged during the journal's editorial and peer review processes and acknowledges the anonymous volunteers on PubPeer and PPS "Feet of Clay" [2] whose observations complemented our internal investigation.

#### References

1. Cabanac G., Labbé C. and Magazinov A., Tortured phrases: A dubious writing style emerging in science. Evidence of critical issues affecting established journals, *ArXiv, abs/2107.06751* (2021).
2. Cabanac G., Chain retraction: how to stop bad science propagating through the literature [Comment], *Nature* 632(8027) (2024), 977-979. <https://doi.org/10.1038/d41586-024-02747-1>

출처: SAGE/Journal of Intelligent & Fuzzy Systems 웹 페이지

#### ● (대응 방안)

- ① 연구자의 최종 검토 원칙 유지: 생성형 AI가 수정한 내용은 반드시 연구자가 직접 검토하여 연구 결과의 의미와 논리적 흐름이 정확하게 유지되는지 확인해야 함.
- ② 전문 용어 및 개념 확인: 논문 수정 과정에서 학문 분야의 전문 용어와 개념이 정확하게 사용되었는지 점검해야 함.
- ③ 저널 정책 확인: 논문 제출 전에 해당 학술지의 생성형 AI 사용 정책을 확인하고 필요한 경우 AI 활용 사실을 명시해야 함.

# 대학 연구자를 위한 생성형 AI 연구윤리 가이드



# 생성형 AI 활용 후 인용

- 1) 생성형 AI 활용 사실의 공개
- 2) 저널별 정책 확인하기

## 제5장

## 제5장

### 생성형 AI 활용 후 인용

#### 1) 생성형 AI 활용 사실의 공개

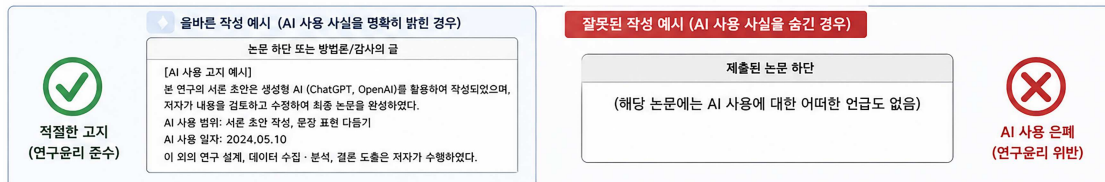
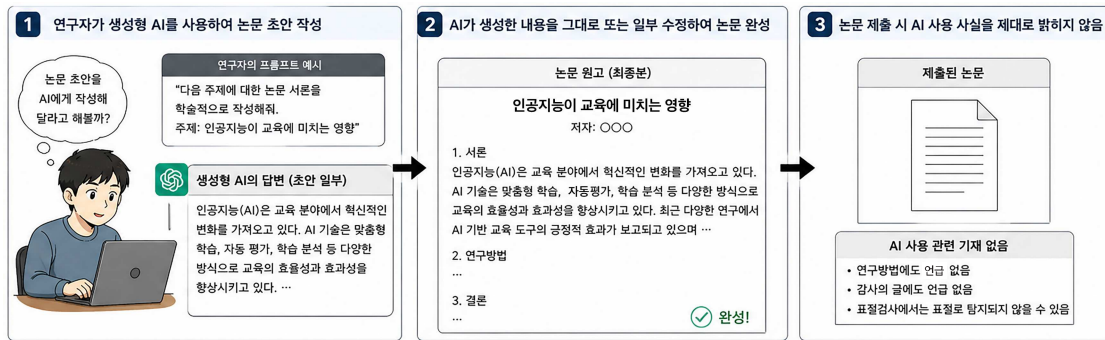
- (정의) 연구자는 논문 작성이나 연구 과정에서 생성형 AI를 활용한 경우 그 사용 사실과 활용 범위를 투명하게 공개할 필요가 있음.

- 생성형 AI는 학술 자료와 동일한 의미의 참고문헌으로 인용되는 것이 아니라 연구 과정에서 활용된 도구(tool)로서 명시되는 것이 일반적임.

#### 예시

#### 생성형 AI 활용 사실 공개 관련 문제점

생성형 AI를 사용하여 논문을 작성했음에도 이를 제대로 밝히지 않는 경우, 연구부정행위로 이어질 수 있음.



출처: "생성형 AI를 사용하여 논문 작성에 도움을 받았음에도 이를 밝히지 않은 것은 문제가 될 수 있음을 보여주는 이미지 만들어줘" ChatGPT 5.4, 2026.05.01.



## 예시

## 생성형 AI 활용 사실 명시

- 생성형 AI 사용 후 생성형 AI 활용 여부를 표시해야 함.
- 학문 분야별 특성과 학술지별 정책 등에 맞게 아래의 예시로 보여준 내주, 각주, 참고문헌, 활용 공개서 중 선택하여 활용 가능함.

□ 내주 및 각주(예시): 내주에 생성형 AI 활용을 표시한 후, 각주로 생성형 AI 활용 여부 및 내용을 공개할 수 있음.

~This study, while providing valuable insights into the role of collaborative communication, government message authenticity, and susceptibility to interpersonal influence in driving ESG-related behaviors in public institutions, is not without limitations (Lee & Hwang 2026, Open AI).<sup>1)</sup>

1) “해당 문장의 문법적 오류를 수정해줘” ChatGPT(OpenAI, 5.3ver.2026), 2026.05.03. 생성형 AI가 제시한 응답을 연구자가 검토 및 수정하여 본문에 반영하였음.

□ 참고문헌(APA 7판 형식)(예시): 내주와 각주에 생성형 AI 활용 여부를 표시하였다면, 참고문헌을 표시해야 함.

OpenAI. (2026). ChatGPT (GPT-5.3 version) [Large language model].  
<https://chat.openai.com/>

□ 논문 섹션별로 생성형 AI 활용여부를 아래와 같이 표시할 수 있음.

Acknowledgment(사사)에서의 예시

“Generative AI was used to assist with the formatting and grammar of this manuscript(May. 3rd, 2026. ChatGPT5.2)”

Methodology Section에서의 예시

“Text was generated with the assistance of Open AI’s GPT-4(OpenAI, 2023) for initial drafting of the literature review section. The final text was extensively edited and approved by all authors(May. 3rd, 2026. ChatGPT5.2)”

□ AI 활용 공개서(예시): 연구 수행 과정에서 활용된 생성형 AI 도구의 사용 내역과 활용 범위를 공개하기 위한 예시로서, 연구 수행 과정에서 생성형 AI 활용의 적절성·책임성·투명성 등을 검토하기 위한 참고 자료로 활용할 수 있음.

#### 1. 기본 정보

| 항목        | 내용                         |
|-----------|----------------------------|
| 논문(연구과제명) | 한국과 터키의 민주주의 비교 연구         |
| 저자(연구책임자) | 홍길동                        |
| 연구수행일자    | 2026. 3. 1. ~ 2026. 5. 10. |

#### 2. 활용 도구 정보 및 내용

| 항목          | 내용                           |                                                                                                                                                                                                                                                             |
|-------------|------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 도구명<br>(버전) | Chat GPT(5.2)                |                                                                                                                                                                                                                                                             |
| 활용 내용       | 선행 연구<br>탐색 및<br>선행 연구<br>정리 | <ul style="list-style-type: none"> <li>연구 주제와 관련된 키워드 도출 및 초기 문헌 매핑 단계에서 ChatGPT를 활용하여 검색어 후보와 관련 이론적 틀을 탐색하였다.</li> <li>단, 실제 인용된 모든 문헌은 연구자가 Google Scholar, Web of Science, RISS를 통해 직접 검색·확인하였으며, AI가 생성한 서지정보는 일체 인용에 사용하지 않았다.</li> </ul>             |
|             | 번역<br>및<br>윤문                | <ul style="list-style-type: none"> <li>영문 초록 및 본문 일부의 번역 초안 작성에 DeepL Pro와 ChatGPT를 사용하였으며, 한국어 본문의 어색한 표현을 다듬는 윤문(language editing) 과정에서 Claude를 활용하였다.</li> <li>모든 AI 생성 텍스트는 연구자가 검토 및 수정하여 반영하였다.</li> </ul>                                            |
| 활용 기간       | 2026. 3. 1. ~ 2026. 5. 10.   |                                                                                                                                                                                                                                                             |
| 도구명<br>(버전) | Claude(2025)                 |                                                                                                                                                                                                                                                             |
| 활용 내용       | 통계<br>분석                     | <ul style="list-style-type: none"> <li>R 코드 작성 및 디버깅 과정에서 Claude를 활용하였다. 구체적으로 다층모형(multilevel modeling) 분석을 위한 lme4 패키지 사용법 확인, 오류 메시지 해석, 시각화 코드(ggplot2) 작성 보조에 활용하였다.</li> <li>분석 설계, 변수 선정, 결과 해석은 연구자가 독립적으로 수행하였으며, 모든 분석 결과는 연구자가 검증하였다.</li> </ul> |
| 활용 기간       | 2026. 3. 1. ~ 2026. 5. 10.   |                                                                                                                                                                                                                                                             |

#### 3. 최종 책임 확인

본 연구의 아이디어, 자료, 분석, 해석, 결론 및 최종 원고에 대한 책임은 저자(연구자)에게 있으며, 저자는 본 연구가 재현 가능하도록 생성형 AI 활용 과정을 투명하고 성실하게 관리하였으며, 활용 내용이 위와 같음을 확인합니다.

날짜: \_\_\_\_\_ 서명: \_\_\_\_\_

## 2) 저널별 정책 확인하기

- (정의) 최근 많은 학술지와 출판사는 생성형 AI 활용과 관련된 별도의 정책을 제시하고 있음. 연구자는 논문 제출 전에 해당 학술지의 생성형 AI 사용 정책을 확인하고 이를 준수해야 함.
- (저널 공통) 저널별로 다양한 정책을 사용하고 있으나, 아래 사항들은 대부분의 저널에서 공통적으로 포함하고 있음.
  - ① AI는 연구 내용에 책임을 질 수 없기에 저자가 될 수 없음.
  - ② AI 사용 후 최종 책임은 연구자에게 있음.
  - ③ AI가 출력한 결과는 출처를 검증하고, 내용의 정확성과 편향성 여부를 검토해야 함.
  - ④ AI를 사용한 경우, 이를 공개해야 함.
- (저널별 정책 사례) 현재 국내외 저널에서는 AI 사용과 관련하여, 다양한 정책을 수립하여 안내하고 있음.
  - ① 생성형 AI 도구를 활용하였을 경우, 논문 내에 저자 주, 각주 등에 명시해야 한다고 언급함.

### 한국영화학회 영화연구 윤리규정(개정 2026.04.28.)

#### 제6조 (인공지능(AI) 기술 활용 윤리)

1. 저자는 ChatGPT, Claude, Gemini 등 생성형 AI 도구를 연구 및 논문 작성 과정에서 활용하였을 경우 이를 논문 내에 명시하여야 한다.
2. AI 활용 명시는 다음 방법 중 하나 이상을 사용한다.
  - 가) 저자 주(author's note) 또는 감사의 글(acknowledgements) 섹션에 사용한 AI 도구의 명칭, 버전 및 활용 목적을 서술
  - 나) 해당 부분의 각주에 AI 도구 활용 사실 명시
3. 생성형 AI는 연구의 저자(author)로 인정되지 않는다. AI 생성 콘텐츠를 자신의 독창적 결과물로 제시하는 것은 연구부정행위에 해당한다.
4. AI 도구를 활용하여 생성된 텍스트, 이미지, 번역물 등의 저작권 및 윤리적 고려사항은 다음과 같다.
  - 가) AI 생성 콘텐츠에 포함된 제3자 저작물에 대한 저작권 침해 문제는 저자가 책임진다.
  - 나) AI 생성 콘텐츠의 사실 정확성 및 학술적 타당성 검증 책임은 저자에게 있다.
  - 다) AI를 활용한 영상 분석, 자막 생성, 이미지 인식 등의 경우 사용한 도구와 방법론을 연구방법 섹션에 기술하여야 한다.
5. AI 활용 연구의 투명성 확보를 위해 다음 사항을 연구방법 또는 부록에 기재하도록 권고한다.
  - 가) 사용한 AI 도구의 명칭 및 버전
  - 나) 주요 프롬프트 또는 지시어(prompt) 내용
  - 다) AI 결과물의 검증 및 수정 과정

② AI-assisted editing tool(보조 교정 도구)로만 사용한 경우에는 별도 공개가 필요 없다고 언급함.

한국유통학회 연구윤리 규정(개정 2025.07.03.)

제13조(인공지능(AI) 사용에 관한 보고)

저자는 생성형 AI(Generative AI) 및 AI 보조 도구(AI-assisted tools)를 사용한 경우, 이를 원고 투고 시 논문 제목에 각주로 표시하여 보고하여야 한다. 다만, AI 보조 교정(AI-assisted editing)을 목적으로 활용한 경우에는 이를 별도로 공개할 필요가 없다. 여기서 AI 보조 교정이란, 저자가 작성한 원고의 가독성과 문체를 개선하고, 문법, 철자, 문장부호, 어조 및 참고문헌의 오류를 수정하는 작업을 의미한다. 그 외 작업(생성적인 편집작업, 자율적인 콘텐츠 생성, 데이터 분석, 연구에 대한 통찰 도출 등)에 사용한 경우 이를 명시하는 별도의 선언문을 추가하여야 한다.

선언문: 본 논문의 작성 과정에서 저자(들)는 [사용한 도구/서비스 명]을(를) 사용하여 [사용 목적]을 수행하였습니다. 도구/서비스를 사용한 후, 저자(들)는 모든 내용을 검토 및 수정하였으며, 출판된 논문의 내용에 대한 전적인 책임을 집니다.

Statement: In the process of preparing this work, the author(s) used [NAME of TOOL/SERVICE] in order to [REASON]. After using this tool/service, the author(s) carefully reviewed and edited the content as necessary, and take(s) full responsibility for the final version of the published article.

③ AI 생성 결과물을 독창적 성과로 허위 표시하거나, 그대로 제출하는 경우 연구부정행위에 해당 될 수 있다고 언급함.

문학과영상학회 생성형 AI 활용에 관한 연구윤리규정(시행 2025.09.01.)

제3조(부정행위 금지 원칙)

생성형 AI의 생성 결과물을 자신의 독창적 연구 성과로 허위 표시하거나 생성 내용을 그대로 제출하는 것은 부정행위(표절, 위조, 변조 등)에 해당할 수 있다.

한국종교학회 생성형 AI 활용 연구윤리규정(제정 2025.12.29.)

제3조(생성형 AI 활용의 기본 원칙)

2. 생성형 AI의 결과물을 자신의 독창적 연구 성과로 허위 표시하거나, 생성 내용을 그대로 제출·게재하는 행위는 표절·위조·변조 등 연구부정행위에 해당할 수 있다.

한국영어영문학회 생성형 AI 활용에 관한 연구윤리규정(적용 2025.10.01.)

제3조(생성형 AI 활용 원칙)

2. 생성형 AI의 생성 결과물을 자신의 독창적 연구 성과로 허위표시 하거나 생성 내용을 그대로 제출하는 것은 부정행위(표절, 위조, 변조 등)에 해당할 수 있다.

- ④ AI로 생성한 이미지 등은 원칙적으로 논문에 제시하는 것을 허용하지 않고 있으며, AI가 부적절하게 사용되었다고 판단하면 심사를 진행하지 않음.

#### Science 저널의 AI 활용 정책

AI는 저자가 될 수 없고, AI 사용 내역은 공개해야 하며, AI 생성 이미지 사용은 명시적 허락없이 허용되지 않음.

##### Artificial intelligence (AI)

AI-assisted technologies [such as large language models (LLMs), chatbots, and image creators] do not meet the Science journals' criteria for authorship and therefore may not be listed as authors or coauthors, nor may sources cited in Science journal content be authored or coauthored by AI tools. Authors who use AI-assisted technologies as components of their research study or as aids in the writing or presentation of the manuscript should note this in the cover letter and in the methods section or acknowledgments section of the manuscript (the location and details of the declaration will depend on use case; see our AI usage guidelines for specifics). Authors are accountable for the accuracy of the work and for ensuring that there is no plagiarism. They must also ensure that all sources are appropriately cited and should carefully review the work to guard against bias that may be introduced by AI. Editors may decline to move forward with manuscripts if AI is used inappropriately. Reviewers may not enter any part of the manuscript into an LLM or other AI tool because this could breach the confidentiality of the manuscript. However, reviewers may use AI tools to revise their own writing, provided that the system does not save or use inputs to train the model and that such use is transparently declared in the review. Reviewers are solely accountable for outputs created with the help of generative AI tools. See our AI usage guidelines.

**AI-generated images and other multimedia are not permitted in the Science journals without explicit permission from the editors.** Exceptions may be granted in certain situations—e.g., for images and/or videos in manuscripts specifically about AI and/or machine learning. Such exceptions will be evaluated on a case-by-case basis and should be disclosed at the time of submission. The Science journals recognize that this area is rapidly developing, and our position on AI-generated multimedia may change with the evolution of copyright law and industry standards on ethical use.

##### Images

Images presented in research papers should correctly represent the original data. No part of a digital image may be selectively manipulated or altered. When figures are assembled from multiple images or non-contiguous portions of the same image, a line or space should indicate the border between the separate parts. Science journal editors may use Profig to screen images.

출처: <https://www.science.org/content/page/science-journals-editorial-policies#authorship>

### Springer Nature의 AI 활용 정책

AI는 저자가 될 수 없고, AI 생성 이미지 포함 허용 안 함, 원고를 AI에 업로드하는 것은 금지됨.

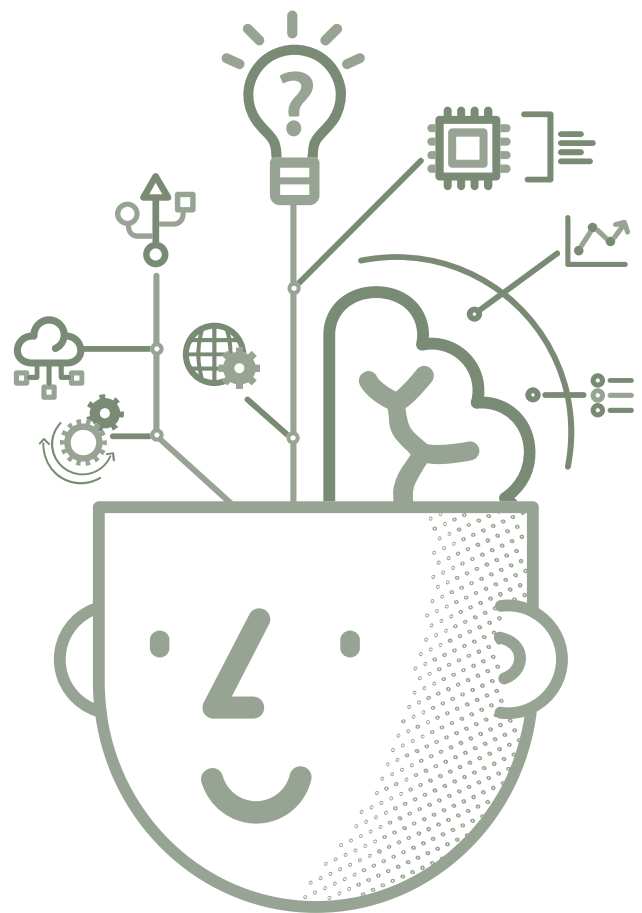
#### Artificial Intelligence

We believe in using AI responsibly and for the benefit of the research community, our authors, editors and readers, and our staff.

We achieve this by being committed to adopting an ethically-focused approach while using, designing, developing, and deploying AI based solutions. We design and use solutions which contain AI or are enabled by AI responsibly, making sure that we consider and mitigate any negative impact, be it societal or environmental. We place human-centered values at the heart of our approach to the responsible use of AI, and these are reflected in our AI Principles and editorial policies. Below are links to our current editorial policies concerning the use of AI. Springer Nature is monitoring ongoing developments in this area closely and will review and update these policies as appropriate.

- ✔ SN does not attribute authorship to AI.
- ✔ SN does not allow the inclusion of generative AI images in our publications.
- ✔ SN asks peer reviewers not to upload manuscripts into generative AI tools.

출처: <https://www.springernature.com/gp/policies/editorial-policies>



# 대학 연구자를 위한 생성형 AI 연구윤리 가이드

---

# 연구자를 위한 생성형 AI 활용 체크리스트

## 제6장

## 제6장

### 연구자를 위한 생성형 AI 활용 체크리스트

본 체크리스트를 참고하여 생성형 AI 활용 과정에서 필요한 사항을 점검하고, 책임 있는 연구 활동에 적극 활용해 주시기 바랍니다.

| 단계별      | 연번 | 질문                                                                                                    | 예 | 아니오 | 해당 없음 |
|----------|----|-------------------------------------------------------------------------------------------------------|---|-----|-------|
| 연구 설계 단계 | 1  | 생성형 AI를 활용하여 선행 연구를 탐색한 경우, 주요 학술 데이터베이스(Scopus, Web of Science, Google Scholar 등)를 통해 문헌을 다시 확인하였는가? |   |     |       |
|          | 2  | 생성형 AI가 제시한 연구 동향이나 연구 공백이 실제 선행 연구와 일치하는지 검토하였는가?                                                    |   |     |       |
|          | 3  | 생성형 AI가 제시한 연구 문제나 가설을 그대로 사용하지 않고 선행연구와 이론적 근거를 통해 검증하였는가?                                           |   |     |       |
|          | 4  | 생성형 AI가 제시한 가설 외에 대안 가설이나 경쟁 가설을 함께 검토하였는가?                                                           |   |     |       |
| 연구 수행 단계 | 1  | 연구 데이터(설문 응답, 인터뷰 자료, 개인 정보 등)를 생성형 AI 시스템에 입력할 때, 비식별화를 하였는가?                                        |   |     |       |
|          | 2  | 연구 데이터의 수집, 분석, 처리 과정에서 생성형 AI를 활용한 경우 그 과정이 연구자가 이해하고 검증할 수 있는 방식으로 이루어졌는가?                          |   |     |       |
|          | 3  | 생성형 AI가 제시한 분석 방법이나 코드가 연구 목적과 데이터 특성에 적합한지 검토하였는가?                                                   |   |     |       |
|          | 4  | 데이터 분석 결과가 실제 데이터에 근거하여 해석되었는지 확인하였는가?                                                                |   |     |       |

| 단계별          | 연번 | 질문                                                         | 예 | 아니오 | 해당 없음 |
|--------------|----|------------------------------------------------------------|---|-----|-------|
| 연구 결과 해석 단계  | 1  | 생성형 AI가 분석한 통계 결과와 설명을 검토하고 검증하였는가?                        |   |     |       |
|              | 2  | 생성형 AI가 제시한 결과 해석을 그대로 사용하지 않고 연구자가 직접 검토하고, 필요하다면 수정하였는가? |   |     |       |
|              | 3  | 생성형 AI가 설명한 연구 결과의 해석이 실제 데이터가 지지하는 범위를 넘어 과장되지 않았는가?      |   |     |       |
|              | 4  | 연구 설계의 한계와 데이터의 제약을 함께 제시하였는가?                             |   |     |       |
| 연구 보고 단계     | 1  | 생성형 AI가 생성한 문장이나 내용을 검토하였는가?                               |   |     |       |
|              | 2  | 생성형 AI가 작성한 문장이나 표현을 학문 분야의 전문 용어와 개념에 맞게 수정하였는가?          |   |     |       |
|              | 3  | 참고문헌 목록에 포함된 문헌이 실제 존재하는 문헌인지 확인하였는가?                      |   |     |       |
|              | 4  | 생성형 AI가 제시한 참고문헌의 원문을 확인하였는가?                              |   |     |       |
| 생성형 AI 활용 공개 | 1  | 생성형 AI를 활용한 경우 그 활용 범위를 논문에 명시하였는가?                        |   |     |       |
|              | 2  | 논문 제출 전에 해당 학술지의 생성형 AI 활용 정책을 확인하였는가?                     |   |     |       |
|              | 3  | 생성형 AI는 연구 보조 도구로 활용되었으며 논문의 최종 책임은 연구자에게 있음을 인지하고 있는가?    |   |     |       |
|              | 4  | 연구 결과에 영향을 미친 주요 프롬프트와 응답 내용을 PDF, 링크, 아카이브 등의 형태로 보존하였는가? |   |     |       |

# 대학 연구자를 위한 생성형 AI 연구윤리 가이드

---

## 생성형 AI 활용을 위한 책임 있는 연구 수행 10가지 준칙

### 1 AI를 저자로 표시하지 말라.

- AI는 연구의 책임 주체가 될 수 없으므로 저자 자격을 가질 수 없다.

### 2 AI 사용에 대한 최종 책임은 연구자가 져라.

- 연구 아이디어, 가설 검증, 데이터 해석, 결론 도출의 주체는 항상 연구자이며, 최종 책임도 연구자에게 있다.

### 3 AI가 생성한 정보는 반드시 검증하라.

- 생성형 AI는 틀린 정보(데이터, 이미지, 참고문헌, DOI 등), 존재하지 않는 논문, 잘못된 통계를 제시할 수 있으므로 반드시 검토 후 사용해야 한다.

### 4 개인정보·기밀정보를 AI에 입력하지 말라.

- 설문 응답, 인터뷰 자료, 미공개 연구데이터, 특허·내부자료는 생성형 AI에 입력해서는 안 된다. 데이터를 입력하기 전에 반드시 비식별화하여야 한다.

### 5 AI의 분석 결과를 이해하지 못하면 사용하지 말라.

- 통계, 코드, 수치해석, 시뮬레이션 결과는 연구자가 직접 이해·재현·검증할 수 있을 때만 활용해야 한다.

### 6 AI의 해석을 그대로 믿지 말고 연구자가 직접 판단하라.

- 생성형 AI는 결과를 과장 또는 왜곡하거나 맥락 없이 해석할 수 있으므로, 연구자가 직접 검토하고 판단해야 한다.

### 7 AI 사용 가능 범위를 사전에 확인하라.

- 학술지, 연구기관, 과제 지원기관 등 생성형 AI 사용 기준이 다를 수 있으므로, 논문 및 보고서 작성 전에 허용 범위와 기준을 확인하라.

### 8 AI 사용 사실은 투명하게 공개하라.

- 논문, 보고서, 과제 수행 시 AI를 어디에 어떻게 활용했는지 각주·감사의 글·별도 양식 등으로 밝히는 것이 바람직하다.

### 9 AI 사용 과정을 기록하라.

- 어떤 AI 도구를 사용했고, 어떤 목적으로 사용했는지 등을 기록해 두는 것이 바람직하다.

### 10 연구윤리 기준을 준수하라.

- AI를 사용했더라도 위조, 변조 등 연구부정행위는 그대로 적용된다.

# 대학 연구자를 위한 생성형 AI 연구윤리 가이드

---

- [1] Bommasani, R., Hudson, D. A., Adeli, E., Altman, R., Arora, S., von Arx, S., ... & Liang, P. (2021). On the opportunities and risks of foundation models. arXiv preprint arXiv:2108.07258.
- [2] Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., ... & Wright, R. (2023). "So what if ChatGPT wrote it?" Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- [3] van Dis, E. A. M., Bollen, J., Zuidema, W., van Rooij, R., & Bockting, C. L. (2023). ChatGPT: five priorities for research. *Nature*, 614(7947), 224-226. <https://doi.org/10.1038/d41586-023-00288-7>
- [4] European Commission, Directorate-General for Research and Innovation. (2025). Living guidelines on the responsible use of generative AI in research (2nd ed.). European Research Area Forum. [https://research-and-innovation.ec.europa.eu/document/download/2b6cf7e5-36ac-41cb-aab5-0d32050143dc\\_en?filename=ec\\_rtd\\_ai-guidelines.pdf](https://research-and-innovation.ec.europa.eu/document/download/2b6cf7e5-36ac-41cb-aab5-0d32050143dc_en?filename=ec_rtd_ai-guidelines.pdf)
- [5] International Committee of Medical Journal Editors. (2023). Recommendations for the conduct, reporting, editing, and publication of scholarly work in medical journals. <https://www.icmje.org/recommendations/>
- [6] Creswell, J. W., & Creswell, J. D. (2018). *Research design: Qualitative, quantitative, and mixed methods approaches* (5th ed.). SAGE Publications.
- [7] Saunders, M. N. K., Lewis, P., & Thornhill, A. (2019). *Research methods for business students* (8th ed.). Pearson Education.
- [8] Ji, Z., Lee, N., Frieske, R., Yu, T., Su, D., Xu, Y., Ishii, E., Bang, Y. J., Madotto, A., & Fung, P. (2023). Survey of hallucination in natural language generation. *ACM Computing Surveys*, 55(12), Article 248. <https://doi.org/10.1145/3571730>
- [9] Kerr, N. L. (1998). HARKing: Hypothesizing after the results are known. *Personality and Social Psychology Review*, 2(3), 196-217. [https://doi.org/10.1207/s15327957pspr0203\\_4](https://doi.org/10.1207/s15327957pspr0203_4)
- [10] Tversky, A., & Kahneman, D. (1974). Judgment under uncertainty: Heuristics and biases. *Science*, 185(4157), 1124-1131. <https://doi.org/10.1126/science.185.4157.1124>
- [11] Popper, K. R. (1959). *The Logic of Scientific Discovery*. London: Hutchinson.
- [12] Jaime A. Teixeira da Silva, 2024. literature. *Current Issues in Sport Science (CISS)*, 9(1), 010. <https://doi.org/10.36950/2024.9ciss010>.
- [13] Lee, H. 2025. Culture as a Barrier to Technological Innovation: The Paradox of Monozukuri in Japanese Corporate Governance, *정치정보연구*, 28(2):68-93.

# 대학 연구자를 위한 생성형 AI 연구윤리 가이드

저 자 이호빈 대학연구윤리협의회/충남대학교

감 수 황은성 서울시립대학교  
전주홍 서울대학교

발행·인쇄일자 2026년 6월 30일

발 행 처 한국연구재단

문 의 처 한국연구재단 윤리정책법무팀

편 집 · 제 작 에코디자인

이 책자의 저작권은 재단법인 한국연구재단에 있으며, 무단으로 수정 및 배포를 금지합니다.  
이 책자의 내용을 한국연구재단의 허락없이 영리 목적으로 판매 등의 행위는 저작권법 등  
관련 법에 따라 처벌될 수 있으며, 내용을 인용 시 출처를 밝혀야 합니다.



교육부



한국연구재단



KUCRE  
대학연구윤리협의회