

From Policing to Partnership: Leveraging GenAI for Compassionate and Culturally Responsive Academic Integrity

Hebatallah Shoukry¹, Gule Saman (s.gule@hw.ac.uk)¹, Nidhal Abdulaziz¹ and Clare Thomson²

¹School of Engineering & Physical Sciences, ²Learning and Teaching Academy, Heriot-Watt University

Introduction

This handout explains the rationale, methodology, results, and implications of the pilot study investigating Generative AI (GenAI) as a collaborative partner in academic misconduct decision making.

Study Background

The study was conducted with the Academic Misconduct Committee (AMC) in the School of Engineering and Physical Sciences at Heriot-Watt University. It examines whether GenAI can help committees reach decisions that are fairer, faster, more consistent, more empathetic, and more culturally responsive particularly for a student population with wide-ranging cultural and educational backgrounds.

Research Purpose

Rather than treating AI solely as a source of academic misconduct, the project explores whether AI can support ethical leadership by contributing evidence synthesis, policy alignment, consistency checking, and bias awareness while leaving final authority to human decision-makers.

Methodology

Committee members reviewed three hypothetical misconduct scenarios. For each case, participants first made an independent decision on guilt and penalty. They then reviewed a policy-informed GenAI recommendation for the same case and rated it against five criteria, using Likert-scale questions covering:

- Empathy and context
- Trust and consistency
- Prior information and bias

- Cultural fairness
- Willingness to adopt

Seven AMC members took part in the pilot (n = 7); as a small, single-institution sample, the findings below should be read as indicative rather than generalisable, see Limitations.

Case Studies

Three scenarios were designed to probe different points on the misconduct spectrum, testing both the standardisation benefits of GenAI input and the risks of bias or insensitivity to exceptional circumstances:

- **Case 1 — GenAI-Assisted Coursework Submission** (Year 2 student)
- **Case 2 — Programmatic Misconduct and Code Generation** (Year 4 student)
- **Case 3 — Unauthorised Aid in an Assessment** (Final Year Project student)

Results

Independent academic judgements varied considerably by case. In Case 1, 42.9% of respondents (3 of 7) judged the student guilty; in Case 2, 85.7% (6 of 7); and in Case 3, 100% (7 of 7). Recommended sanctions also became progressively more severe as the perceived seriousness of the misconduct increased — summarised in Table 1.

Interpretation of Findings

Participants consistently valued human judgement for its ability to weigh context, hardship, and mitigating circumstances factors that are difficult to encode into a policy-informed AI recommendation. At the same time, the high internal consistency of GenAI's recommendations increased participants' trust in the tool, and many supported using GenAI to accelerate decisions in routine, lower-stakes cases.

Ethical Considerations

Respondents raised concerns about cultural fairness, data privacy, transparency of how recommendations are generated, and the risk of GenAI amplifying existing biases in policy or precedent. Given these concerns, participants favoured a human-in-the-loop model, in which GenAI provides a recommendation and supporting rationale, but the committee retains full responsibility and authority for the final decision.

Implications for Higher Education

The findings suggest that GenAI could become an effective decision-support tool within academic integrity systems. Potential advantages include consistency, efficiency, policy

alignment, and structured evidence review. However, human oversight remains essential to maintain fairness, empathy, and accountability.

Limitations and Future Work

- Small, single-institution sample (n = 7) from one committee and one school; results may not generalise across disciplines or institutions.
- Scenarios were hypothetical rather than live cases, which may not fully capture the emotional and procedural weight of a real hearing.
- Planned next steps include a larger cross-institutional sample, live (anonymised) case shadowing, and a longitudinal look at how trust in GenAI recommendations changes with sustained use.

Conclusion

The poster supports a partnership model for AI in academic misconduct decision making. The proposed approach combines the analytical strengths of GenAI with the contextual understanding and ethical judgement provided by academics, creating a more transparent and compassionate decision-making framework.

References

1. B. Moya, S. E. Eaton, H. Pethrick, K. A. Hayden, R. Brennan, J. Wiens, B. McDermott, and J. Lesage, "Academic integrity and artificial intelligence in higher education contexts: A rapid scoping review protocol." *Canadian Perspectives on Academic Integrity*, vol. 5, no. 2, pp. 59-75, 2023.
2. C. K. Y. Chan, "A comprehensive AI policy education framework for university teaching and learning", *International Journal of Educational Technology in Higher Education*, vol. 20, no. 1, pp. 38, 2023. DOI:[10.1186/s41239-023-00408-3](https://doi.org/10.1186/s41239-023-00408-3)
3. S. A. D. Popenici, and S. Kerr, "Exploring the impact of artificial intelligence on higher education", Springer Nature, 2024. DOI: [10.1186/s41039-017-0062-8](https://doi.org/10.1186/s41039-017-0062-8)