



OPEN ACCESS

EDITED BY

Kelly Merrill Jr.,
University of Cincinnati, United States

REVIEWED BY

Carla Freire,
University of Minho, Portugal
Cesare Giulio Ardito,
The University of Manchester,
United Kingdom
Ruth Baker-Gardner,
University of the West Indies, Mona, Jamaica

*CORRESPONDENCE

Tiffany Petricini
✉ tzt106@psu.edu

RECEIVED 12 December 2024

ACCEPTED 11 April 2025

PUBLISHED 01 May 2025

CITATION

Petricini T, Zipf S and Wu C (2025)
RESEARCH-AI: Communicating academic
honesty: teacher messages and student
perceptions about generative AI.
Front. Commun. 10:1544430.
doi: 10.3389/fcomm.2025.1544430

COPYRIGHT

© 2025 Petricini, Zipf and Wu. This is an
open-access article distributed under the
terms of the [Creative Commons Attribution
License \(CC BY\)](#). The use, distribution or
reproduction in other forums is permitted,
provided the original author(s) and the
copyright owner(s) are credited and that the
original publication in this journal is cited, in
accordance with accepted academic
practice. No use, distribution or reproduction
is permitted which does not comply with
these terms.

RESEARCH-AI: Communicating academic honesty: teacher messages and student perceptions about generative AI

Tiffany Petricini^{1*}, Sarah Zipf² and Chuhao Wu²

¹Department of Communication and Media, Penn State Erie, The Behrend College, Erie, PA, United States, ²Teaching and Learning with Technology, The Pennsylvania State University (PSU), University Park, PA, United States

Integrating generative AI (GenAI) in higher education presents both opportunities and challenges, particularly in maintaining academic integrity. This study explores faculty communication strategies when addressing suspected GenAI misuse, analyzing responses through Gallant's rule-based and integrity-focused frameworks. Data from a survey reveal a dominant reliance on punitive, rule-based approaches, highlighting tensions between students and instructors. While some faculty adopt communicative and educational strategies, fostering trust and collaboration, others exhibit a dismissive stance due to perceived administrative burdens. The findings support the growing research that institutions prioritize educational interventions and support faculty in developing trust-based, proactive strategies for integrating GenAI responsibly.

KEYWORDS

generative AI, academic integrity, instructor communication, higher education, faculty student trust

Introduction

The rapid rise of generative artificial intelligence (GenAI) has transformed and will continue to transform higher education, presenting both opportunities and challenges for faculty and students alike. Beyond mastering how to use these tools effectively, educators must also guide students in ethical and effective GenAI use. Unlike previous waves of educational technology, the speed of GenAI's evolution and the large breadth of applications have created a unique sense of urgency. Institutional policy has not been able to keep up with the rapid development and evolution of these tools. EDUCAUSE's *2024 AI Landscape Study* found that only a small percentage of institutions have established comprehensive AI policies, and most recommend that individual faculty develop course-specific guidelines (Robert, 2024). These policies often address concerns about academic integrity, ethical AI use, and the appropriate application of AI tools in student work (Robert, 2024). While many policies focus on academic integrity and ethical AI use, they often fail to address how faculty communicate expectations, concerns, and trust in student engagement with GenAI.

This study addresses a critical gap in the literature by examining how faculty approach conversations with students about suspected GenAI misuse. Prior research has explored the ethical concerns and institutional policies surrounding AI in education (Gallant, 2008; Kumar et al., 2024), yet little attention has been given to the communicative strategies faculty employ in these interactions. Given that instructor communication plays a crucial role in shaping student perceptions of fairness, trust, and learning outcomes (Eaton, 2021), understanding these dynamics is vital.

Academic integrity, or the honesty with which academic material has been produced, has been a significant concern among higher education for over a century and “has routinely been called on in times of perceived crises in postsecondary education” (Gallant, 2008, p. 2). GenAI has certainly created a crisis as instructors and institution scramble to adapt to its rapid adoption and the challenges. Work in communication studies promotes significant pre-crisis work to yield the best outcomes when crises emerge (Coombs, 2007; Ulmer et al., 2007). However, with the rapid onset of these tools, it stands to reason that reactive approaches are more likely to happen than proactive approaches to curtail the pre-crisis.

Research indicates that while institutional leadership recognize the transformative potential of GenAI, many feel unprepared to support faculty and students (Watson and Lee, 2025). AI literacy is critical to all disciplines and institutions moving forward, although a significant barrier to building AI literacy are the negative misperceptions that students and faculty alike hold about AI (Petricini et al., 2024; Zipf et al., 2024). AI literacy efforts are complicated by students’ fear of being unfairly accused of using AI tools when they haven’t (Zipf et al., 2025). Given past research showing that negative perceptions and misconceptions about generative AI can significantly impact both faculty and students (Petricini et al., 2024; Zipf et al., 2024), decisions faculty make about AI tools and academic integrity in their courses could potentially be fear-based rather than evidence-based.

Studies show more students will cheat because AI technology makes it easier to do so, even when students do not trust the technology (Dahl and Waltzer, 2024; Robinson and Glanzer, 2017). Recent research has critically evaluated the effectiveness of AI-generated text detection tools. A study by Weber-Wulff et al. (2023) assessed 14 tools, including popular platforms like Turnitin and GPTZero, and found that all scored below 80% in accuracy (Weber-Wulff et al., 2023). These tools often misclassify human-written content as AI-generated (false positives) and fail to identify AI-generated text (false negatives). Additionally, simple manipulations of AI-generated content, such as paraphrasing, can significantly reduce the accuracy of these detectors (Perkins et al., 2024). The overall evidence suggests that automated detection tools are currently unreliable for distinguishing between human and AI-generated texts.

Eaton (2021) reports students of minority populations are more likely to be blamed for cheating and to suffer from unintentional harm than those students who fit the homogeneous culture. While AI technologies can support and help several student populations, including nonnative English speakers, Eaton (2021) warns about the potential collateral damage that can be done when policing or punishing academic integrity violations. This intersection of academic integrity, equity, and GenAI is an emerging area of research, specifically around the communication efforts of instructors.

Understanding how instructors communicate with students regarding possible academic integrity violations is paramount, as the ways in which instructors communicate regarding issues can significantly impact students’ lives and well-being. In 1972, researchers asked over 150 students to identify two of the most negative experiences in their lives (Branan, 1972). Distressingly, teachers were most often cited as the source of these experiences, with reported behaviors including “humiliation in front of a class, unfairness in evaluation, destroying self-confidence, personality conflicts, and embarrassment” (p. 82). With the possibility of having such profound

impacts on student well-being, understanding how instructors communicate about academic integrity violations is paramount to student success as we move forward with integrating GenAI tools into education.

Academic integrity policies in higher education often fall into two broad categories: rule-based and integrity-focused approaches (Gallant, 2008). Rule-based approaches prioritize adhering to strict institutional policies, often relying on punitive measures, like academic sanctions and AI detection tools to enforce compliance. In contrast, integrity-focused approaches emphasize fostering students’ ethical decision-making, prioritizing conversations and trust-building over punishment. Gallant (2008) argues that an overemphasis on rule enforcement can create adversarial relationships between faculty and students, whereas an integrity-focused model helps students internalize academic honesty as a core value. To help guide practices and investigate the above concerns, this study explored communication strategies by faculty, either planned or past, when they suspected students of academic integrity violations that used GenAI by applying Gallant’s framework. By categorizing faculty responses as rule-based, integrity-focused, or emerging categories (dismissive and collaborative), this research highlights how faculty responses shape student perceptions of fairness, trust, and institutional integrity expectations.

This study addresses the following research questions:

1. How do faculty members communicate with students about suspected GenAI misuse in academic settings?
2. To what extent do faculty responses align with Gallant’s (2008) rule-based and integrity-focused framework?
3. What additional patterns (beyond rule-based and integrity-based approaches) emerge in faculty communication strategies?
4. What institutional and pedagogical factors influence faculty members’ decision-making in cases of suspected GenAI misuse?

Examining these varied approaches and their implications, this study contributes to the growing discourse on AI literacy, academic integrity, and instructor-student trust in higher education. The findings emphasize the need for institutions to move beyond reactive, punitive measures and instead support faculty in adopting proactive, trust-building communication strategies that promote responsible AI use.

Materials and methods

Survey method

We utilized a mixed method approach, combining both quantitative and qualitative data, collected through survey method. As we are combining both numerical and textual data to answer the research question, a mixed method approach is appropriate for this study (Tashakkori and Creswell, 2007). Building on a survey from the prior year (Petricini et al., 2024), we updated language and added additional items based on changes with GenAI in higher education. The self-administered online questionnaire included 21 items and two open-ended prompts. A portion of the survey containing 12 items is presented in this study related to faculty and student perceptions of AI and faculty’s open-ended responses to one open-ended prompt related to faculty’s approach to students’ improper use of GenAI tools.

This study was conducted at a large, research-intensive university located in the mid-Atlantic region of the United States. The university has over 40,000 students and more than 7,000 faculty members. After approval from the institutional review board, we sent emails to administrative and academic leadership to send to their respective faculty, staff, and student rosters, soliciting volunteers for our study. The emails included a link to a Qualtrics questionnaire and used implied consent with minimal risk to participants. The online questionnaire was available from March to May 2024.

Data analysis

The online questionnaire items were analyzed with descriptive statistics, finding the means, standard deviation, and t-tests between faculty and student responses. The data from the open-ended prompt were analyzed at the semantic level and thematically and within the conceptual framework. Responses were coded grounded in [Gallant \(2008\)](#) framework that differentiates between rule-based and integrity approaches. Two researchers independently analyzed and coded all responses and met to review. The initial interrater reliability testing showed a moderate level of agreement ($k = 0.56$). The researchers discussed the differences and recoded the items two more times, reaching an interrater reliability of a very high level of agreement ($k = 0.89$).

During coding, two more codes emerged. One was termed as “dismissive,” characterized by faculty responses that either downplayed the need for direct communication about GenAI usage or assumed that assignment instructions were sufficient in conveying expectations. The other addition was “collaboration,” in which responses were neither rules-based, integrity focused, nor dismissive. Instead, these answers indicated approaches in which instructors prioritized communication and problem-solving together with the student.

Survey results

In total, 85 faculty and 66 students submitted responses to the survey with 69 faculty responding to the open-ended question about their approach to inappropriate GenAI usage. The demographic distribution of faculty and student are illustrated in [Table 1](#).

Results from the selected survey items ([Table 2](#)) show students self-reported their experience with GenAI as higher than faculty (4.03 and 3.36 respectively). Interestingly, students agreed their instructors’ policies as clear (4.03) though disagreed that they have confidence in their instructors’ use of GenAI in the classroom (2.77) and neutral that instructors are responsible with AI use (3.0). It is unclear what makes students think policies are clear, when they lack confidence in instructors’ use of GenAI. Instructors are less sure about their own policies (3.59) but strongly believe students should be taught how to use GenAI responsibly (4.61).

Distrust between instructors and students with GenAI tools is a significant insight from this research and remains unchanged from the previously conducted study ([Petricini et al., 2024](#)). One area of distrust stems from transparency: while instructors only somewhat agreed they should tell students about their use of GenAI tools (3.43), students felt transparency was important (4.42). Students are worried about being falsely blamed for using GenAI (4.25) and feel mostly neutral that their instructors would believe them if they said they had not used GenAI tools (3.34).

Instructors think students misuse GenAI (3.78) and disagree that AI detection tools are accurate (2.74). Both students and

TABLE 1 Survey respondents’ demographics.

Demographics	Category	Faculty	Student
Gender	Woman	48.20%	57.60%
	Man	42.40%	33.30%
	Other	9.40%	9.10%
Ethnicity	White	78.80%	78.80%
	Black/African	2.40%	3.00%
	Asian	1.20%	1.50%
	Hispanic/LatinX	1.20%	0.00%
	Native American	0.00%	0.00%
	Other	16.50%	16.70%
Discipline	STEM	22.50%	48.90%
	Applied disciplines	28.10%	38.30%
	Arts	5.60%	8.50%
	Social sciences	33.40%	2.10%
	Other	10.30%	2.10%

faculty somewhat agree that using GenAI is a violation of academic integrity (4.33 and 4.06 respectively) but are less agreed that the use of GenAI violates institutional policy (3.35 and 3.72, respectively), a potential indication that institutional policies are unclear ([Table 2](#)).

Open-ended response results

We received a total of sixty-nine responses to our open-ended prompt: How have you or how would you approach a student you suspect of cheating or plagiarizing with generative AI?

The findings indicate significant variation in how faculty approach academic integrity concerns related to GenAI. Using [Gallant \(2008\)](#) framework, we categorized faculty responses into rule-based, integrity-focused, dismissive, and collaborative approaches ([Table 3](#)).

- **Rule-based approaches (39.1%):** The dominant approach, where faculty strictly adhered to policy guidelines and often defaulted to punitive measures such as failing students or reporting them for misconduct. Example responses included:

“If I suspect genAI use, I immediately follow the institution’s academic integrity process. The student receives a failing grade.”

- **Integrity-focused approaches (24.6%):** Faculty using this approach engaged students in discussions about AI ethics and responsible usage. One instructor described their approach as:

“I would have a conversation with the student to understand why they used AI and discuss how to use it ethically moving forward.”

- **Collaborative approaches (20.3%) (New Category):** Some faculty viewed AI-related concerns as a chance to engage students in ethical decision-making and course policy development, rather than strictly enforcing rules.

TABLE 2 Selected items from survey.

Item	Mean ± Standard deviation		T-test	p-value
	Faculty	Student		
I'm experienced with using generative AI.	3.36 ± 1.18	4.03 ± 1.14	−3.37	0.001
Academic integrity issues are a risk of using generative AI.	4.33 ± 0.88	4.06 ± 1.18	1.52	0.131
I/My instructor have a clear policy about students' use of GenAI in class.	3.59 ± 1.19	4.03 ± 1.14	−2.20	0.029
Using generative AI tools to complete coursework violates academic integrity policies at the University.	3.35 ± 1.14	3.72 ± 1.28	−1.74	0.084
AI detection tools are accurate (i.e., Turnitin, GPTZero).	2.74 ± 1.01			
I worry that my instructor may blame me for using generative AI on assignments or exams even if I did not.		4.25 ± 0.95		
My instructor will trust me when I say I did not use generative AI.		3.34 ± 1.14		
I am confident in my instructors' use of generative AI in the classroom.		2.77 ± 0.97		
My instructors practice generative AI use responsibly in the classroom.		3 ± 1.11		
Students misuse generative AI.	3.78 ± 1.00			
Students should be informed when instructors use generative AI for instructional support (i.e., developing exam questions, creating writing prompts, grading, etc.).	3.43 ± 1.16	4.42 ± 0.81	−5.90	<0.001
We need to teach students more about using generative AI responsibly.	4.61 ± 0.77			

TABLE 3 Frequency count of codes.

Code	Frequency by percentage (n)
Rule	39.1% (27)
Integrity	24.6% (17)
Collaborative	20.3% (14)
Dismissive	8.7% (6)
Did not answer	7.2% (5)

“Instead of banning AI, I work with students to decide when and how it should be used in my course.”

- **Dismissive approaches (8.7%) (New Category):** Some faculty members indicated they would not actively pursue AI-related academic integrity violations, either due to skepticism about detection tools or institutional burdens.

“The process is too cumbersome. I won’t report students unless it’s blatant.”

These findings suggest that Gallant’s framework is useful but incomplete in capturing faculty responses to AI-related integrity concerns. The emergence of dismissive and collaborative categories reflects the complexity of faculty perspectives, which are influenced not only by pedagogical philosophy but also by institutional policies and workload constraints.

Discussion

Our findings indicate that faculty overwhelmingly rely on punitive, rule-based approaches when addressing suspected GenAI

misuse, aligning with prior research that academic integrity policies often emphasize compliance over education (Gallant, 2008; Kumar et al., 2024). This tendency may stem from faculty uncertainty about GenAI, as our quantitative data reveal, and that instructors report significantly less confidence in their own AI literacy compared to students. The uncertainty likely contributes to a reliance on instructional policies rather than proactive engagement with students.

Also, the study uncovers a mismatch between student and faculty perceptions of AI policies and instructor competence. While students report that policies are “clear,” they simultaneously express low confidence in their instructors’ AI-related decisions. The contradiction suggests that policies may be well-articulated but fail to establish trust or provide meaningful guidance. This distrust coupled with fears of false accusations may exacerbate inequities (Zipf et al., 2025), and resistance to AI literacy initiatives.

A minority of faculty adopt a collaborative, trust-based approach, treating AI-related academic integrity concerns as opportunities for learning rather than as infractions requiring punishment. These instructors engage students in dialog and emphasize ethical AI use, aligning with Eaton (2021) argument that institutions should prioritize integrity-building over surveillance and punitive enforcement. However, the limited number of faculty employing this approach makes it clear that institutions must provide more training and support to help faculty shift from reactive to proactive strategies.

Finally, a small subset of faculty adopt a dismissive stance, avoiding engagement with GenAI-related academic integrity issues, sometimes due to perceived administrative burdens. Similarly, Kumar et al.’s (2024) conclude that faculty are often overwhelmed by enforcement, leading to disengagement. The institutional response to AI use in education must address this burden.

Practical recommendations

We seek practical recommendations for faculty and institutions and offer the following recommendations. First, a shift must occur from punitive to educational approaches. Institutions should develop clearer guidelines that encourage an integrity-based approach rather than defaulting to rule enforcement. Faculty training should emphasize how to communicate academic integrity expectations transparently and constructively. Second, AI literacy must be enhanced among faculty. Workshops, learning communities and institutional support should focus on equipping instructors with AI literacy skills. Faculty need support integrating AI ethically into coursework, moving beyond a binary framing of AI as either permissible or prohibited. Third, transparent and fair academic policies are warranted. Institutions should create uniform processes for handling AI-related integrity cases to reduce perceived unfairness. Clear guidelines need to be communicated to students, with explicit guidance about what constitutes responsible AI use rather than assuming all AI use is misconduct. Similarly, institutions need to consider how faculty's use of AI might be made transparent to students and each other, so that students' learning and assessment are not entirely outsourced to an AI tool. In doing so, the institution will remain focused on maintaining the quality and authenticity of students' educational experience. Last, encourage trust-based, dialogic communication. Faculty should be encouraged to engage in open conversations with students about their AI use rather than defaulting to accusations. Institutions can support this process by providing training and mediation resources and modeling communication practices based on mutual respect and transparency.

Conclusion

This study highlights several insights into the evolving landscape of academic integrity as GenAI becomes a central element in higher education. Our findings demonstrate that while many faculty members default to rule-based, punitive communicative approaches when addressing suspected AI misuse, a significant potential for more constructive, trust-building strategies exists.

These integrity-focused and communicative approaches not only align with Gallant's framework but also create opportunities for fostering deeper understanding and collaboration between students and instructors.

By addressing the gap between faculty and student perceptions of AI policies and their implementation, this research forefronts the importance of transparency and proactive communication. Faculty development programs must prioritize AI literacy and the cultivation of educational strategies that encourage student growth and responsibility rather than fear-based compliance. Additionally, institutions need to address faculty workload and provide clearer guidelines to reduce the perceived administrative burdens that often lead to dismissive attitudes.

Looking ahead, future research should explore the long-term impact of integrity-focused strategies on student outcomes. How diverse student populations experience and respond to AI-related academic integrity policies should also be investigated, particularly in

the context of equity and inclusion. As GenAI tools continue to advance, creating dynamic and ethical educational environments will require ongoing collaboration between researchers, educators, and institutional leaders.

Data availability statement

The datasets presented in this article are not readily available because of participant consent and institutional approval. Requests to access the datasets should be directed to TP, tzt106@psu.edu.

Ethics statement

This study was approved by the Pennsylvania State University Institutional Review Board, study #00021869. Written informed consent from the participants or the participants' legal guardian/next of kin was not required to participate in this study in accordance with the national legislation and the institutional requirements.

Author contributions

TP: Writing – original draft, Writing – review & editing. SZ: Writing – original draft, Writing – review & editing. CW: Writing – original draft, Writing – review & editing.

Funding

The author(s) declare that no financial support was received for the research and/or publication of this article.

Conflict of interest

The authors declare that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.

Generative AI statement

The authors declare that Gen AI was used in the creation of this manuscript. Generative AI was used to proofread, to brainstorm, and to do external non-included analyses.

Publisher's note

All claims expressed in this article are solely those of the authors and do not necessarily represent those of their affiliated organizations, or those of the publisher, the editors and the reviewers. Any product that may be evaluated in this article, or claim that may be made by its manufacturer, is not guaranteed or endorsed by the publisher.

References

- Branan, J. M. (1972). Negative human interaction. *J. Couns. Psychol.* 19, 81–82. doi: 10.1037/h0032030
- Coombs, W. T. (2007). Ongoing crisis communication: Planning, managing, and responding. 2nd Edn. Los Angeles, CA, USA: SAGE Publications.
- Dahl, A., and Waltzer, T. (2024). A canary alive: what cheating reveals about morality and its development. *Hum. Dev.* 68, 6–25. doi: 10.1159/000534638
- Eaton, S. E. (2021). Plagiarism in higher education: Tackling tough topics in academic integrity. Santa Barbara, CA, USA: Libraries Unlimited.
- Gallant, T. B. (2008). Academic integrity in the twenty-first century: a teaching and learning imperative. *ASHE High. Educ. Rep.* 33, 1–143. doi: 10.1002/aehe.3305
- Kumar, R., Eaton, S. E., Mindzak, M., and Morrison, R. (2024). “Academic integrity and artificial intelligence: an overview” in Second handbook of academic integrity, 1583–1596. doi: 10.1007/978-3-031-54144-5_153
- Perkins, M., Roe, J., Vu, B. H., Postma, D., Hickerson, D., McGaughran, J., et al. (2024). Simple techniques to bypass GenAI text detectors: implications for inclusive education. *Int. J. Educ. Technol. High. Educ.* 21:53. doi: 10.1186/s41239-024-00487-w
- Petricini, T., Zipf, S., and Wu, C. (2024). Perceptions about generative AI and ChatGPT use by faculty and students. *Transform. Dialogues Teach. Learn. J.* 17, 63–87. doi: 10.26209/td2024vol17iss21825
- Robert, J. (2024). 2024 EDUCAUSE AI landscape study: EDUCAUSE Publications.
- Robinson, J. A., and Glanzer, P. L. (2017). Building a culture of academic integrity: what students perceive and need. *Coll. Stud. J.* 51, 209–221.
- Tashakkori, A., and Creswell, J. W. (2007). Editorial: exploring the nature of research questions in mixed methods research. *J. Mixed Methods Res.* 1, 207–211. doi: 10.1177/1558689807302814
- Ulmer, R. R., Seeger, M. W., and Sellnow, T. L. (2007). Effective crisis communication: Moving from crisis to opportunity. Thousand Oaks, CA, USA: SAGE Publications.
- Watson, C. E., and Lee, R. (2025). Leading through disruption: Higher education executives assess AI’s impacts on teaching and learning: AAC& U & Elon University.
- Weber-Wulff, D., Anohina-Naumeca, A., Bjelobaba, S., Foltýnek, T., Guerrero-Dib, J., Popoola, O., et al. (2023). Testing of detection tools for AI-generated text. *Int. J. Educ. Integr.* 19, 1–39. doi: 10.1007/s40979-023-00146-z
- Zipf, S., Petricini, T., and Wu, C. (2024). “AI monsters: an application to student and faculty knowledge and perceptions of generative AI” in The role of generative AI in the communication classroom. (Hershey, PA, USA: IGI Global), 284–299.
- Zipf, S., Petricini, T., and Wu, C. (2025). Using the information inequity framework to study GenAI equity: analysis of educational perspectives. *Inf. Res.*