

Generative AI and Academic Integrity: A Sociological Perspective

Teena¹, Megha Sharma², Seema Boliya³

^{1,2,3}Assistant Professor, M.L.V. Government College, Bhilwara

Abstract

Generative artificial intelligence (AI) has entered higher education quickly, and it has brought back old sociological questions about how academic misconduct gets defined, spread, and controlled. This paper does not treat AI-assisted writing as just a quicker way to plagiarise. Instead, it argues that generative AI shakes the very norms that academic integrity has always relied on. Using classical sociological theories of deviance — strain theory, neutralization theory, social learning theory, and institutional theory — along with recent research on AI in higher education, the paper looks at how AI blurs the line between help and authorship, how “cheating” is a socially built category applied unevenly, how detection tools carry linguistic and economic bias, and why institutional policies often look better than they work. The paper concludes that academic integrity is not something technology can simply threaten or destroy. It is something built continuously through institutional design, peer relationships, and resources that are not shared equally among students.

Keywords: generative artificial intelligence, academic integrity, sociology of deviance, higher education, academic dishonesty

Introduction

When ChatGPT became publicly available in late 2022, universities reacted quickly — new proctoring tools, updated honor codes, and syllabus lines banning “unauthorized use of artificial intelligence.” At first glance, generative AI looked like an old problem with a faster engine: plagiarism, just quicker. But treating it only as something to detect and punish misses the more interesting sociological story. Academic integrity was never a fixed moral fact waiting to be enforced. It is a social arrangement built and rebuilt through shared expectations about effort, authorship, and fairness among students, faculty, and institutions (Cabrera Vélez et al., 2026). Generative AI does not just make cheating easier — it exposes how fragile these shared assumptions always were, about where a student’s words end and a machine’s begin, and about what actually counts as legitimate help.

This paper looks at generative AI and academic integrity through a sociological lens, not a purely pedagogical or legal one. Instead of asking “how do we catch AI-assisted cheating?”, it asks “what does the arrival of generative AI reveal about how academic norms are made, contested, and unevenly enforced?” The argument moves in five steps. First, it revisits classical theories of deviance — strain theory, neutralization theory, social learning theory, and institutional theory — to understand why students turn to AI and how they justify it. Second, it looks at how generative AI blurs the line between assistance and authorship. Third, it treats academic dishonesty as something socially constructed, shaped by who gets to define it. Fourth, it examines how AI-related suspicion falls unevenly on students, especially

multilingual and international ones. Finally, it considers how institutions respond — often symbolically rather than practically — and what a more honest approach to integrity could look like.

Theoretical Foundations: Sociological Lenses on Deviance and Norms

Sociologists did not need generative AI to explain why people break rules they claim to value. Robert Merton's strain theory is one of the oldest and still most useful starting points. Merton (1938) argued that deviance often comes not from rejecting cultural goals but from a gap between those goals and the legitimate means available to reach them. In today's university, the goal is clear: good grades, credentials, and the credibility they bring. But the legitimate means — enough study time, language fluency, stable housing, money for tutoring — are not shared equally. Recent research on cheating supports this: students often describe cheating less as rejecting academic values and more as a way to cope with workloads and pressures that outstrip the time and resources they have (Garcines et al., 2024). In Merton's terms, generative AI works as an innovative adaptation — a new means, not a new goal, used by students who still want to succeed on the institution's own terms.

A second lens comes from Gresham Sykes and David Matza's (1957) theory of neutralization, which explains how people who break norms still manage to see themselves as good people. Their five techniques — denial of injury, denial of the victim, condemning the condemners, appeal to higher loyalties, and denial of responsibility — apply to generative AI use almost too well. A student who tells themselves that no one is really hurt by asking a chatbot to tighten a paragraph is using denial of injury. One who argues the assignment was pointless busywork is condemning the condemners. Denial of responsibility takes an interesting shape here, because the tool itself offers an easy way to spread out authorship: it becomes simple to say "I didn't really write that, the AI did," as if the prompt, the choices, and the final submission were not decisions at all.

Albert Bandura's (1986) social learning theory adds something the other two theories miss: the role of peers. Behaviour, in this view, is learned mostly by watching others. Seeing someone use a shortcut and get away with it makes that shortcut more likely to come to mind the next time a person feels similar pressure. What sets generative AI apart from older forms of dishonesty is how fast and how openly this learning happens — in shared documents, group chats, and campus gossip about which instructors do and don't check for AI use. Cheating, seen this way, is rarely a lone moral failure. It spreads through networks of trust, much like any other learned social behaviour.

Finally, institutional theory helps explain why universities have struggled to respond with any real consistency. John Meyer and Brian Rowan (1977) argued that organizations often adopt formal structures not because they work well, but because they signal legitimacy to outsiders — what they called rationalized myth and ceremony. A quickly written AI-use policy, a detection-software subscription, or a new syllabus clause can serve this ceremonial purpose even when everyone involved doubts it will change much. This does not mean such policies are dishonest; it means they exist partly to show that the institution is managing a new risk, not necessarily because they manage it well.

Generative AI and the Erosion of Shared Normative Boundaries

There is also a quieter shift happening beneath the policy debates — one about what even counts as a rule. Academic integrity has always rested on a rough, if imperfect, line between acceptable help and unacceptable substitution. A tutor could explain a concept but not write the essay; a classmate could proofread but not restructure an argument; a parent could encourage but not draft. Generative AI does not

break this line so much as make it hard to find. A student who asks a chatbot to explain a reading and then paraphrases it in their own words is doing something similar to what students have always done with tutors and reference books. A student who asks the same chatbot for a finished paragraph and submits it unchanged is doing something closer to buying an essay. In between lies a large, ordinary, and mostly unmapped middle ground — brainstorming with AI, using it to reorganize a rough draft, or asking it to check an argument's logic.

Recent research confirms this ambiguity is not a temporary confusion that better guidelines will fix — it is a structural feature of the technology itself. A systematic review of forty-one studies on generative AI and academic integrity found that institutions consistently struggle to separate legitimate support from prohibited substitution, partly because the same tool can be used for both within one assignment (Bittle & El-Gayar, 2025). A separate multi-country study reached a similar conclusion: faculty, students, and administrators often hold very different ideas of what “appropriate use” even means, with some instructors treating any AI involvement as misconduct and others welcoming it as digital literacy (Yusuf et al., 2024). Sociologically, this is not just a communication problem to be fixed with clearer wording. It is what happens when a technology arrives faster than the shared norms needed to regulate it — a genuinely Durkheimian moment of anomie, where old rules no longer clearly apply and new ones have not yet formed.

None of this is entirely new. Generative AI complicates the idea of authorship in ways that existed before this moment, but AI now pushes them into the open. Group projects, co-authored papers, and heavily edited dissertations have always involved shared contribution. What generative AI does is compress that sharing into a single, often invisible exchange between one student and one tool — stripped of the visibility that once made collaboration easy to see and govern.

The Social Construction of “Cheating” in the Age of AI

A recurring idea in the sociology of deviance is that misconduct is not a natural category waiting to be discovered. It is defined by whoever has the power to label behaviour that way, and that labelling shifts with institutional interest, historical moment, and who is doing the judging. Academic dishonesty is no different. What seemed like an unremarkable use of a calculator, a spell-checker, or a citation manager a generation ago was, when first introduced, treated with much the same suspicion now aimed at generative AI. The comparison should not be stretched too far — generative AI can produce far more of an assignment's actual content than a spell-checker ever could — but the pattern of alarm followed by slow, uneven acceptance is a familiar one.

Jeffrey Dixon (2025), a sociologist writing on this exact question, makes a related point worth taking seriously: public conversation about generative AI in higher education has focused almost entirely on student cheating, while quietly setting aside harder questions about the responsibility of the technology companies that built these tools and the institutions that adopted them, often without real consultation of the students and faculty who would have to live with the results. Placing the entire burden of ethical use on individual students, he argues, hides a wider distribution of responsibility that a purely disciplinary view conveniently ignores. This is a useful correction to accounts that treat academic dishonesty as simply a matter of individual character. If integrity really is a property of a whole educational ecosystem — including assessment design, institutional transparency, and corporate accountability — then blaming only the student's choices is itself a kind of social construction, one that conveniently reduces the responsibility of more powerful actors.

This is not an argument that using generative AI to complete an assignment dishonestly is acceptable. It is an argument that the category of “cheating” is being actively negotiated in real time — by administrators writing policy, instructors setting their own informal rules, and students testing what will and will not be tolerated. Sociological attention should be paid to who gets to do that negotiating, and on what terms.

Structural Inequality and the Politics of Detection

One troubling finding in recent research is about who actually gets caught, or suspected, of AI-assisted misconduct. AI-text detectors, the tools many institutions use to police this boundary, do not fail randomly. Research on these detectors found they consistently misclassified writing by non-native English speakers as AI-generated far more often than writing by native speakers — a pattern the authors linked to the more formulaic, less varied sentence structures that language learners often use, which happen to resemble patterns typical of machine-generated text (Liang et al., 2023). In practice, this means international and multilingual students, already navigating unfamiliar academic conventions in a second or third language, face a higher risk of false accusation from the very technology meant to protect academic integrity.

This is not a small technical glitch — it is a structural inequity with real disciplinary consequences, and it falls hardest on students who are already on the margins. A social-justice review of generative AI in education takes this further, connecting detector bias to bigger questions about whose labour built these systems, whose data trained them, and whose language norms they were calibrated against (Healy, 2023). The same review points out that a technology trained mostly on text from wealthier, English-dominant parts of the internet is unlikely to treat all forms of academic writing as equally legitimate — a bias that repeats itself first in what the model generates, and then again in what detection tools flag as suspicious.

A similar imbalance runs along economic lines. Students with reliable internet, paid AI subscriptions, and earlier exposure to these tools through well-resourced schools arrive at university already comfortable using generative AI productively and discreetly. Students without that exposure are more likely to use these tools clumsily, in ways more likely to get flagged, or not use them at all — missing out on the efficiency gains their peers enjoy. Seen this way, the panic around AI-assisted cheating risks doing what moral panics usually do: focusing suspicion on the most visible and least powerful people, while leaving the conditions that produced the behaviour, and the tools used to police it, largely unexamined.

Institutional Response: Ceremony, Ecosystems, and the Limits of Policy

Faced with this uncertainty, most universities have reacted the way organizations typically do when facing a new, poorly understood risk: quickly, publicly, and somewhat symbolically. A recent review of regulatory frameworks across South American higher education found that despite widespread institutional anxiety about generative AI, there is almost no binding legal language that actually defines where educational support ends and academic dishonesty begins. Existing rules mostly restate general ethical principles about responsible technology use, without turning them into enforceable standards (Cabrera Vélez et al., 2026). This gap between visible activity and real clarity is exactly what Meyer and Rowan’s (1977) theory of organizations would predict: the policy exists largely to show the institution is taking the problem seriously, whether or not it actually changes what happens in a classroom.

Some scholars have proposed more systemic alternatives to this ceremonial pattern. Instead of treating integrity as a matter of catching individual violations after the fact, an ecosystem-based approach reframes it as something that emerges from how assessments are designed, how policies are communicated, and how institutional culture treats disclosure and trust — arguing that enforcement-first models risk turning

faculty into punitive gatekeepers rather than teachers. This fits with the earlier point about learned behaviour: if cheating spreads through networks the way any social practice does, a lasting response has to work at the level of norms and trust within those networks, not just individual detection.

There is a pattern in where governance frameworks do exist: they tend to be strongest exactly where enforcement is weakest. Peru's national legislation promoting AI adoption, for example, mentions education in general terms but has no real mechanism for verifying academic integrity in practice, leaving the actual work to individual universities that often lack the resources or agreement to do it well (Cabrera Vélez et al., 2026). The same pattern shows up across the wider literature: institutions are quick to state a position on generative AI, and slow — or unable — to build the infrastructure that would make that position meaningful (Bittle & El-Gayar, 2025).

Discussion: Toward a Sociological Reframing

Put these ideas together, and a different picture emerges from the one usually heard on campus: generative AI has not created a new moral crisis in higher education so much as made an old one visible. Academic integrity was never held together by universal moral agreement — it was always a fragile, unevenly enforced set of norms, dependent on reasonably stable lines between help and substitution, on students facing broadly similar pressures and resources, and on institutions being able to claim, plausibly, that they were policing misconduct fairly. Generative AI puts strain on all three at once: it blurs the line between assistance and authorship, it favours students with more resources and prior exposure while penalizing others through biased detection, and it shows how much existing policy has worked symbolically rather than practically.

A sociological reframing along these lines does not mean giving up on academic integrity, nor does it excuse students who use these tools to avoid learning altogether. What it does suggest is that treating integrity violations as purely individual moral failures, fixable with better detection software, misreads the problem. If strain theory is right that innovation follows from a mismatch between goals and means, then addressing that strain — through fairer workloads, clearer expectations, and assessments that don't simply invite substitution — will do more good than any detector. If neutralization theory is right that students need a story in which their conduct isn't really wrong, then treating every student as a suspect risks reinforcing exactly the adversarial mindset that makes those stories easier to tell. And if institutional theory is right that much of current policy is ceremonial, real progress will require institutions to admit their existing frameworks were built for appearances rather than function, and to redesign them accordingly.

Conclusion

Generative AI arrived in higher education faster than any shared understanding could be built to receive it, and the confusion that followed has often been framed as a crisis of student character. A sociological view suggests a different, more demanding diagnosis. Academic integrity is not a fixed thing that AI threatens to steal. It is a social accomplishment, produced continuously through institutional design, peer networks, and structural conditions that are not shared equally across the student population. The strain that pushes students toward AI-assisted shortcuts, the neutralizations that make those shortcuts feel acceptable, the peer networks that normalize them, and the ceremonial policies that fail to really govern them — none of these are new problems invented by large language models. What generative AI has done is compress and intensify processes that sociologists have studied for nearly a century, forcing universities

to face, more urgently than before, the gap between the integrity they claim to protect and the conditions under which that protection is actually possible.

References

1. Bandura, A. (1986). *Social foundations of thought and action: A social cognitive theory*. Prentice-Hall.
2. Bittle, K., & El-Gayar, O. (2025). Generative AI and academic integrity in higher education: A systematic review and research agenda. *Information*, 16(4), Article 296. <https://doi.org/10.3390/info16040296>
3. Cabrera Vélez, J. P., Llanos García, R. V., & Bonilla Fierro, L. F. (2026). Generative AI and the assurance of academic integrity in the context of South American higher education. *Frontiers in Education*, 11, Article 1796556. <https://doi.org/10.3389/feduc.2026.1796556>
4. Dixon, J. C. (2025, November 17). Student cheating dominates talk of generative AI in higher ed, but universities and tech companies face ethical issues too. *The Conversation*. <https://theconversation.com/student-cheating-dominates-talk-of-generative-ai-in-higher-ed-but-universities-and-tech-companies-face-ethical-issues-too-268167>
5. Garcines, K. M. A., Estender, M. R., Epondol, J. M., Uy, S. D., Maldepeña, K. M. V., & Cuevas, J. F., Jr. (2024). Stories behind academic cheating: Cheaters' perspective. *International Journal of Research and Innovation in Social Science*, 8(12), 2013–2024. <https://doi.org/10.47772/IJRISS.2024.8120170>
6. Healy, M. (2023). Approaches to generative artificial intelligence, a social justice perspective. *arXiv*. <https://doi.org/10.48550/arXiv.2309.12331>
7. Liang, W., Yuksekgonul, M., Mao, Y., Wu, E., & Zou, J. (2023). GPT detectors are biased against non-native English writers. *arXiv*. <https://doi.org/10.48550/arXiv.2304.02819>
8. Merton, R. K. (1938). Social structure and anomie. *American Sociological Review*, 3(5), 672–682. <https://doi.org/10.2307/2084686>
9. Meyer, J. W., & Rowan, B. (1977). Institutionalized organizations: Formal structure as myth and ceremony. *American Journal of Sociology*, 83(2), 340–363. <https://doi.org/10.1086/226550>
10. Sykes, G. M., & Matza, D. (1957). Techniques of neutralization: A theory of delinquency. *American Sociological Review*, 22(6), 664–670. <https://doi.org/10.2307/2089195>
11. Yusuf, A., Pervin, N., & Román-González, M. (2024). Generative AI and the future of higher education: A threat to academic integrity or reformation? Evidence from multicultural perspectives. *International Journal of Educational Technology in Higher Education*, 21, Article 21. <https://doi.org/10.1186/s41239-024-00453-6>