

Generative AI and Academic Integrity in Higher Education: Challenges and Future Directions

Dr. G. Purushothaman¹, Dr. S. Ganapathy², Mr. Saurabh Jaiswal^{3*}, Mr. Thanga Kumaran M⁴

^{1,2,4} S. A. College of Arts & Science, University of Madras, India

¹ Assistant Professor & Head, PG Department of Commerce (General)

² Assistant Professor & Head, Department of Corporate Secretaryship

⁴ Assistant Professor & Head, Department of English

³ Assistant Professor, Department of Business Administration,

Prof. Rajendra Singh (Rajju Bhaiya) University, Prayagraj, India

¹ profpurushg@gmail.com, ² drprofganapathy@gmail.com, ³ s.jaiswal84@yahoo.com, ⁴ drthangamphdel@gmail.com

Abstract: The public release of ChatGPT in late 2022, and the wave of generative artificial intelligence (GenAI) tools that followed it, has unsettled long-standing assumptions about how learning is demonstrated and assessed in higher education. Essays, reports, code, and even reflective writing can now be produced in seconds by systems whose output is fluent, personalised, and largely indistinguishable from student work. This paper examines the resulting collision between GenAI and academic integrity from a governance and assurance perspective. Drawing on the rapidly growing literature published since 2022, it maps the principal challenges facing institutions: the unreliability and demonstrated bias of AI-text detection tools, the erosion of assessment validity, widening equity gaps, fragmented and reactive policy, limited faculty capacity, and the contamination of scholarly work by fabricated references and hallucinated content. The paper argues that detection-centred enforcement is a structurally weak control and proposes instead a layered institutional framework in which policy and governance, pedagogy and assessment redesign, and technology-based assurance operate as mutually reinforcing controls, sustained by a continuous audit and improvement cycle. Future directions are discussed, including two-lane assessment models, AI literacy as a graduate attribute, provenance and watermarking infrastructure, and the emergence of academic integrity as an auditable domain of institutional risk management.

Keywords: generative artificial intelligence; academic integrity; higher education; assessment; AI detection; governance; ChatGPT; audit and assurance

I. INTRODUCTION

Academic integrity has never been a static concept. Each generation of educational technology, from the pocket calculator to the search engine to the essay mill, has forced universities to renegotiate what counts as legitimate assistance and what counts as misconduct (Bretag, 2016). Yet the arrival of generative artificial intelligence marks a discontinuity rather than another incremental step. Earlier threats to integrity, such as plagiarism from published sources or contract cheating, involved copying or purchasing text that already existed somewhere, which meant it could in principle be found, matched, and attributed (Foltýnek et al., 2019; Lancaster & Cotarlan, 2021). Large language models produce text that has never existed before. There is no source document to match against, no ghostwriter to trace, and no transaction record to subpoena. The evidentiary foundation on which plagiarism management was built for three decades has simply fallen away.

The speed of the change has compounded its depth. ChatGPT reached an estimated hundred million users within two months of its November 2022 launch, making it one of the fastest-adopted consumer applications in history, and students were among its earliest and most enthusiastic adopters (Dwivedi et al., 2023; Rudolph et al., 2023). Within a single academic term, instructors across disciplines reported receiving submissions they suspected were machine-generated but



could not prove to be so. Institutions responded in strikingly different ways: some banned the tools outright, some rushed to procure detection software, and others encouraged experimentation, often within the same university system (Moorhouse et al., 2023). Three years on, the dust has not settled so much as the landscape has kept moving, with multimodal models, AI agents, and assistants embedded directly inside word processors making the boundary between the student's work and the machine's contribution progressively harder to draw.

This paper takes stock of that unsettled landscape. Its purpose is threefold. First, it synthesises the emerging research literature on GenAI and academic integrity to identify the principal challenges confronting higher education institutions. Second, reflecting the concerns of information governance, control, and audit professionals, it reframes academic integrity in the GenAI era as a problem of institutional risk management and assurance rather than one of individual student morality alone. Third, it sketches future directions, both pedagogical and technological, that appear most promising for sustaining trust in academic credentials. Throughout, we argue that the instinctive institutional response, namely an arms race of detection and evasion, is not merely losing but unwinnable in its current form, and that durable solutions will come from redesigning assessment, building AI literacy, and embedding integrity into governance structures that can be monitored, audited, and improved over time.

The remainder of the paper is organised as follows. Section 2 describes the generative AI landscape as it bears on higher education. Section 3 revisits the concept of academic integrity itself and asks how it must be reinterpreted when machine assistance is ubiquitous. Section 4 analyses six interlocking challenge domains. Section 5 proposes a layered governance and assurance framework, and Section 6 discusses future directions before Section 7 concludes.

II. THE GENERATIVE AI LANDSCAPE IN HIGHER EDUCATION

Generative AI refers to a class of machine learning systems, most prominently large language models (LLMs), that produce novel text, images, code, and other media in response to natural-language prompts. The technical breakthrough underpinning current systems was the demonstration that transformer-based models trained on internet-scale corpora acquire remarkably general capabilities, including the ability to perform tasks they were never explicitly trained on, simply from a handful of examples or instructions (Brown et al., 2020; Bommasani et al., 2021). GPT-4 and its contemporaries extended these capabilities to graduate-level examinations, professional certification tests, and multimodal reasoning over images and documents (OpenAI, 2023). What matters for higher education is less the architecture than the consequence: nearly every genre of academic writing that universities use to certify learning can now be generated on demand, at negligible cost, in seconds.

Empirical benchmarking studies quickly confirmed what instructors suspected. Susnjak (2022) demonstrated that ChatGPT could produce plausible, well-structured answers to open-book online examination questions across disciplines. Yeadon et al. (2023) found that AI-written short-form physics essays earned marks comparable to those of typical students, concluding bluntly that the format was no longer a defensible assessment instrument. A large multi-institutional study in engineering education found that GenAI output passed a substantial share of existing assessments outright and came close in many others, with performance varying by task type rather than by discipline (Nikolic et al., 2023). Similar results have accumulated in law, medicine, business, and computing, and each new model generation has narrowed whatever gaps remained (Lo, 2023; Kasneci et al., 2023).

It is worth stressing how quickly the landscape has continued to shift since those early studies, because institutional policy has consistently lagged behind capability. Figure 1 juxtaposes the major capability milestones with the corresponding integrity responses in higher education. The pattern is one of persistent asymmetry: capability moves in months, policy in years. By the time many universities had drafted their first ChatGPT policies, the technology those policies described had already been superseded by multimodal systems, retrieval-augmented tools that cite real sources, and AI assistance woven invisibly into mainstream word processors, browsers, and integrated development environments (Farrokhnia et al., 2024). The question facing an integrity office in 2026 is no longer whether a student used a chatbot, but which of the half-dozen AI layers between the student's first keystroke and the submitted file counts as legitimate support.



Generative AI capability milestones

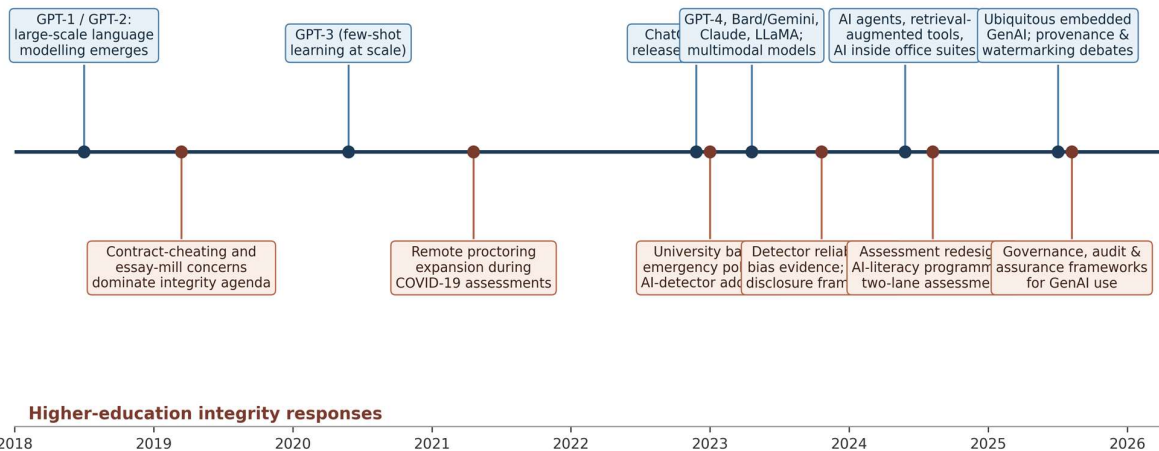


Figure 1. Generative AI capability milestones (top) and higher-education integrity responses (bottom), 2018–2025. The persistent lag between capability and institutional response illustrates the governance asymmetry discussed in Section 4. Source: Authors' synthesis of the literature.

Student adoption, meanwhile, has normalised. Surveys conducted from 2023 onward consistently report that a majority of students use GenAI tools in some form, most commonly for explanation, summarisation, brainstorming, and language polishing, and that students themselves draw nuanced, if inconsistent, distinctions between acceptable assistance and cheating (Chan & Hu, 2023; Sullivan et al., 2023). Notably, students frequently report wanting clearer guidance rather than prohibition; ambiguity, not malice, drives much of what institutions later classify as misconduct. This observation, repeated across national contexts, is central to the governance argument developed later in this paper: where the rules are unclear, unclear behaviour follows.

III. RETHINKING ACADEMIC INTEGRITY IN THE AGE OF GENERATIVE AI

The International Center for Academic Integrity defines integrity through six fundamental values: honesty, trust, fairness, respect, responsibility, and courage (International Center for Academic Integrity, 2021). Nothing in that definition mentions technology, and deliberately so; integrity is a property of conduct, not of tools. Yet the operational machinery through which universities have enforced integrity, similarity scores, originality declarations, invigilated examinations, was built around a specific threat model in which dishonesty meant appropriating existing text or another person's labour. GenAI breaks that threat model without necessarily breaking the underlying values, and much of the current confusion stems from conflating the two.

Eaton (2023) captures the shift with the notion of a “postplagiarism” era, in which hybrid human–AI writing becomes the norm rather than the exception, and in which the historically central question, did you write these words yourself, loses much of its diagnostic power. If a student drafts an argument, asks a model to restructure it, edits the result, and verifies every claim, the words are not wholly theirs, but the intellectual work arguably is. Conversely, a student may type every word personally while outsourcing all of the thinking to a model through iterative prompting. Word-level authorship and cognitive ownership have come apart, and integrity frameworks anchored to the former now misclassify in both directions (Perkins, 2023).

A more defensible framing treats academic integrity as honesty about process: transparent disclosure of what assistance was used, how, and for what. This mirrors long-standing scholarly norms, where researchers acknowledge statistical



consultants, language editors, and funding sources without those acknowledgements diminishing authorship. Several scholars have accordingly proposed disclosure scales or permitted-use tiers that specify, at the level of an individual assessment task, whether AI use is prohibited, permitted with acknowledgement, or actively required (Perkins, 2023; Chan, 2023). The reframing matters for enforcement too. A rule that says “do not use AI” is unverifiable and therefore unenforceable in unsupervised settings; a rule that says “disclose your AI use accurately” converts the integrity question into one of honest reporting, which can be corroborated through process evidence such as drafts, version histories, and oral defence. From a control standpoint, the second rule is auditable where the first is not, a distinction that professionals in assurance roles will recognise immediately.

None of this implies that all AI use is benign. Wholesale outsourcing of assessed work to a machine defeats the certifying purpose of assessment just as surely as contract cheating does, and institutions are entitled to prohibit it. The point is rather that prohibition must be targeted, articulable, and matched to assessment design, because blanket bans that cannot be enforced corrode respect for the rules that can be (Cotton et al., 2024; Dawson, 2021).

IV. CHALLENGES

The challenges GenAI poses to academic integrity are frequently discussed as a single problem, “AI cheating,” but they are better understood as six interlocking domains, each with a distinct risk profile and a distinct set of feasible controls. This section examines each in turn.

4.1 The Detection Dilemma

The first institutional reflex after ChatGPT's release was technological: if machine-generated text is the problem, machine-detected text must be the solution. That reflex has not survived contact with the evidence. In the most comprehensive independent evaluation to date, Weber-Wulff et al. (2023) tested fourteen detection tools, including Turnitin's AI-writing indicator and GPTZero, and found that none achieved both high accuracy and high reliability; performance degraded sharply when AI text was lightly paraphrased, translated, or interleaved with human writing, and the tools produced troubling rates of both false negatives and false positives. Theoretical work reinforces the empirical picture: Sadasivan et al. (2023) showed that as language models improve, the statistical distance between human and machine text shrinks toward zero, meaning that reliable detection may be impossible in principle for state-of-the-art models, and that simple paraphrasing attacks defeat even watermark-based schemes under realistic conditions.

The false-positive problem deserves particular emphasis because its costs fall on the innocent. Liang et al. (2023) demonstrated that widely used detectors systematically misclassify text written by non-native English speakers as AI-generated, apparently because such writing exhibits the lower lexical variability the detectors treat as a machine signature. For internationally diverse institutions, this is not a marginal edge case; it means the control itself discriminates against a protected group of students. A university that initiates misconduct proceedings on the strength of a detector score alone is therefore acting on evidence that is unreliable, biased, and unexplainable, a combination that would fail any serious evidentiary standard and has already led several institutions to disable detector scores entirely. The defensible role for detection tools, if any, is as one weak signal among many, never as the basis for an allegation (Weber-Wulff et al., 2023; Dawson, 2021).

There is also a second-order harm: the arms-race dynamic. Every advance in detection prompts counter-tools, “humanisers” and paraphrasers marketed explicitly for evading detectors, which shifts student effort from learning to laundering and normalises an adversarial relationship between students and institutions (Lancaster, 2023). Controls that train the population to defeat them are poor controls.

4.2 Assessment Validity and Security

Detection failures would matter less if assessments themselves were robust, but much of the assessment estate in higher education was designed for a pre-GenAI world. Dawson (2021) usefully distinguishes assessment security, the ability to assure that work was produced under the claimed conditions, from assessment validity, the ability of the task to measure the intended learning. GenAI attacks both simultaneously. Take-home essays, online quizzes, discussion posts, standard



programming exercises, and reflective journals all now score poorly on security, because a capable model can complete them, and consequently on validity, because a passing grade no longer evidences the intended capability (Susnjak, 2022; Nikolic et al., 2023). Swiecki et al. (2022) argue that this moment exposes weaknesses that predate GenAI: many traditional assessments were always proxies of convenience, measuring the artefact rather than the learning, and the technology has merely made the gap between proxy and construct impossible to ignore.

The tempting response, moving everything back into invigilated examination halls, carries its own validity costs, narrowing what can be assessed to what can be performed under time pressure without resources, which resembles few authentic professional contexts. Remote proctoring, the pandemic-era alternative, brings well-documented privacy, equity, and reliability problems of its own (Dawson, 2021). Assessment design in the GenAI era therefore faces a genuine trilemma among security, validity, and scalability, and Section 6 returns to the emerging compromise positions.

4.3 Equity, Bias, and Fairness

GenAI redistributes advantage among students in ways that integrity policy has barely begun to acknowledge. Access is stratified: frontier models sit behind subscriptions, and students who can pay obtain measurably stronger assistance than those using free tiers, layering a new digital divide onto existing socioeconomic ones (Kasneci et al., 2023; UNESCO, 2023). Enforcement is stratified too, as the detector-bias findings discussed above show, with non-native English speakers facing elevated false-accusation risk (Liang et al., 2023). And the underlying models themselves encode biases of their training corpora, meaning students who rely on them uncritically may absorb and reproduce skewed representations, an issue extensively documented in the broader literature on large language models (Bender et al., 2021). Fairness, one of the six fundamental values of academic integrity, is thus implicated on three axes at once: access, accusation, and content. Institutional responses that address only the misconduct dimension while ignoring the equity dimension will resolve neither (Crawford et al., 2023).

4.4 Policy Vacuum and Governance Gaps

Analyses of institutional policy in the period following ChatGPT's release paint a picture of fragmentation. Moorhouse et al. (2023), examining the world's top-ranked universities, found that only a minority had published assessment-specific GenAI guidelines within the first half of 2023, and that existing guidance varied widely in permissiveness, specificity, and enforceability. Many institutions delegated decisions to individual instructors, producing the situation students report as most corrosive: the same behaviour treated as encouraged scholarship in one course and expellable misconduct in the next (Chan & Hu, 2023). Policy has also tended to be reactive and tool-specific, naming particular products rather than defining principles, and thereby ageing almost immediately (Chan, 2023). From a governance perspective, the deeper gap is structural: academic integrity typically sits with academic affairs, AI procurement with information technology, data protection with legal counsel, and assessment design with individual faculties, with no single owner of GenAI risk across the institution. Risks that cross organisational silos and lack an owner are precisely the risks that mature governance frameworks exist to capture, and their absence in this domain is conspicuous (Dwivedi et al., 2023).

4.5 Faculty Capacity and Institutional Readiness

Policies are executed by people, and the people in question, faculty members, report feeling underprepared. Studies of instructor perspectives describe a workforce asked simultaneously to redesign assessments, adjudicate ambiguous integrity cases, model ethical AI use, and keep pace with tools that change quarterly, generally without workload relief, training, or clear institutional backing (Kasneci et al., 2023; Rudolph et al., 2023). The consequences are predictable. Some instructors default to suspicion, over-relying on detector scores with the evidentiary problems already described; others disengage, tacitly accepting AI-produced work rather than pursuing unwinnable cases. Both responses undermine the consistency on which perceived fairness, and therefore voluntary compliance, depends (Bertram Gallant, 2017). Zawacki-Richter et al. (2019) observed even before the GenAI wave that AI-in-education research was proceeding largely without educators; the integrity crisis has now made that absence operationally costly. Building faculty capacity, through



training, exemplars, shared assessment banks, and protected redesign time, is less glamorous than procuring detection software, but the evidence suggests it is considerably more effective (Holmes & Tuomi, 2022).

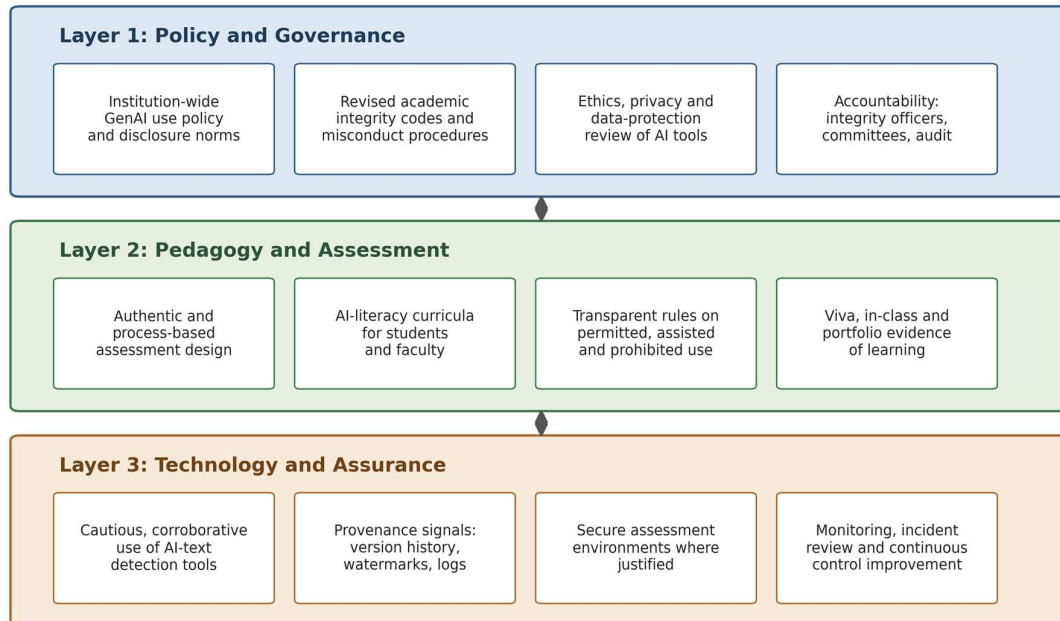
4.6 Hallucination, Fabricated Citations, and Research Integrity

Finally, GenAI threatens integrity even where no one intends to cheat. Language models are prone to hallucination, the confident generation of plausible but false content, including fabricated references complete with real journal names, plausible author lists, and invented DOIs (Ji et al., 2023). Students who use models as search engines can submit work contaminated with non-existent sources in good faith; postgraduate students and faculty face the same risk in literature reviews and grant applications, and journal editors now routinely screen for citation fabrication. The scholarly record itself becomes the victim. This challenge is qualitatively different from deliberate misconduct because the appropriate control is not deterrence but verification competence: teaching users that generative output is a draft to be checked against primary sources, never an authority to be cited, and building verification checkpoints into supervision and editorial workflows (Dwivedi et al., 2023; Farrokhnia et al., 2024). It also underlines a recurring theme of this paper: many GenAI integrity failures are capability failures dressed up as character failures.

V. INSTITUTIONAL RESPONSES: A GOVERNANCE AND ASSURANCE FRAMEWORK

If detection-centred enforcement is structurally weak, what should replace it? The position developed here, consistent with the direction of recent scholarship, is that academic integrity in the GenAI era is best treated as an institutional risk domain requiring layered, mutually reinforcing controls, in the same way organisations treat information security or financial reporting integrity (Chan, 2023; Eaton, 2023). No single control is adequate; the objective is defence in depth, with each layer compensating for the known weaknesses of the others. Figure 2 presents the proposed framework, comprising three layers: policy and governance, pedagogy and assessment, and technology and assurance.

Institutional Governance Framework for Generative AI and Academic Integrity



Layers operate as mutually reinforcing controls: policy legitimises pedagogy; pedagogy reduces reliance on technical policing; technology provides assurance evidence back to governance.

Figure 2. A layered institutional governance framework for generative AI and academic integrity. Source: Authors' conceptualisation.



The policy and governance layer establishes legitimacy and accountability. Its core artefact is an institution-wide GenAI policy that defines principles rather than products: tiers of permitted use, a universal disclosure requirement, and clear allocation of responsibility for keeping guidance current. Misconduct procedures require parallel revision so that allegations rest on corroborated evidence rather than detector scores, and so that honest disclosure is never punished more harshly than concealment, a perverse incentive several early policies inadvertently created (Perkins, 2023; Moorhouse et al., 2023). Crucially, this layer assigns ownership: a named committee or officer responsible for GenAI integrity risk across academic, technological, and legal silos, reporting to senior leadership with the same regularity as other enterprise risks.

The pedagogy and assessment layer is where most of the real risk reduction occurs, because it shrinks the attack surface instead of policing it. Its elements include authentic and process-based assessment design, discussed further in Section 6; explicit AI-literacy instruction for students and staff, covering capabilities, limitations, hallucination, bias, and citation verification; and task-level transparency, so that every assessment states plainly which uses of AI are prohibited, permitted, or expected (Chan, 2023; Sullivan et al., 2023). The literature is consistent that students comply far more readily with specific, reasoned rules than with blanket prohibitions they perceive as arbitrary (Chan & Hu, 2023; Bertram Gallant, 2017).

The technology and assurance layer, finally, is deliberately positioned as the supporting layer rather than the lead. It encompasses the cautious, corroborative use of detection signals; provenance infrastructure such as document version histories, keystroke or draft evidence where proportionate, and emerging watermarking standards; secure assessment environments for the subset of tasks that genuinely require them; and, importantly for the audit community, logging and monitoring sufficient to evidence that the controls in the other two layers are actually operating (Kirchenbauer et al., 2023; Lancaster, 2023; Dawson, 2021). Table 1 summarises how the six challenge domains of Section 4 map onto controls across the three layers.

Table 1. Mapping challenge domains to layered controls (L1 = policy and governance; L2 = pedagogy and assessment; L3 = technology and assurance).

Challenge domain	Principal risk	Primary controls (framework layer in brackets)
Detection dilemma	Unreliable, biased detector evidence; false accusations; arms-race dynamics	Detector output as corroborative signal only (L3); evidence standards in misconduct procedure (L1); process-based evidence such as drafts and vivas (L2)
Assessment validity and security	Existing tasks completable by GenAI; grades no longer certify learning	Authentic, process-based and programmatic assessment redesign (L2); secure environments for high-stakes tasks (L3); assessment risk mapping (L1)
Equity, bias and fairness	Paid-tier advantage; detector bias against non-native writers; biased model content	Institutionally licensed AI access (L1); AI-literacy including bias awareness (L2); validation and bias review of any deployed tools (L3)
Policy vacuum and governance gaps	Fragmented, tool-specific, unowned policy; inconsistent treatment across courses	Principles-based institution-wide policy with named ownership (L1); task-level transparency requirements (L2); compliance monitoring (L3)



Faculty capacity	Inconsistent enforcement; over-reliance on detectors; disengagement	Training, exemplars and redesign time (L2); workload recognition in policy (L1); shared tooling and guidance (L3)
Hallucination and fabricated citations	Good-faith submission of false content; contamination of scholarly record	Verification competence in AI-literacy curricula (L2); citation-verification checkpoints in supervision and editorial workflow (L1/L3)

Static controls, however well designed, decay quickly in a domain where the underlying technology changes quarterly. The framework therefore requires a temporal dimension: a continuous assurance cycle through which risks are re-identified, controls redesigned, implementation verified, incidents investigated fairly, and the whole apparatus independently reviewed. Figure 3 depicts this cycle. Readers from the audit profession will recognise its family resemblance to established plan-do-check-act and IT governance lifecycles, and that resemblance is the point: academic integrity is becoming an auditable domain, and institutions that already possess mature internal audit functions can extend them to this territory at modest marginal cost. Internal audit involvement also supplies something the current debate badly lacks, namely independent evidence about whether integrity controls work, as opposed to whether they exist on paper (Crawford et al., 2023).

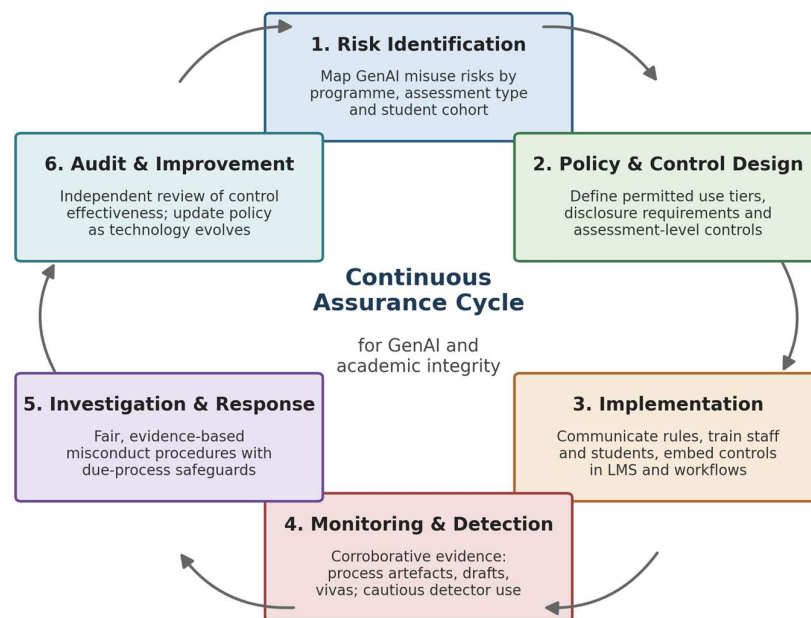


Figure 3. A continuous assurance cycle for generative AI and academic integrity. Source: Authors' conceptualisation.



VI. FUTURE DIRECTIONS

6.1 Two-Lane and Programmatic Assessment

The most consequential developments over the next several years will occur in assessment design. An emerging consensus, sometimes described as a two-lane model, distinguishes assessments whose purpose is secure certification of individual capability from assessments whose purpose is learning through realistic, resource-rich work. In the first lane, a deliberately small set of tasks, vivas, supervised performances, interactive oral assessments, and invigilated demonstrations, is hardened against all forms of outsourcing, machine or human. In the second lane, GenAI use is permitted or expected, disclosure is required, and marking criteria reward judgement, verification, and critical use of the tools rather than unaided production (Dawson, 2021; Swiecki et al., 2022). Programmatic assessment thinking, in which certification decisions rest on triangulated evidence across a whole programme rather than on any single artefact, provides the architecture for combining the lanes. Early adopters report a further benefit: because the second lane mirrors contemporary professional practice, its authenticity improves engagement even as its integrity exposure falls (Swiecki et al., 2022; Nikolic et al., 2023).

6.2 AI Literacy as a Graduate Attribute

A second direction reframes the problem as curriculum rather than compliance. If graduates will work alongside generative systems throughout their careers, then the capacity to use such systems effectively, critically, and honestly is itself a learning outcome, not a threat to learning outcomes. Frameworks such as Chan's (2023) AI Ecological Education Policy Framework embed AI literacy across pedagogical, governance, and operational dimensions, and UNESCO's (2023) global guidance similarly positions human-centred AI competence as a core educational goal. Concretely, this means teaching prompt construction alongside its limits, hallucination and verification, bias recognition, disclosure norms, and the intellectual-property and privacy dimensions of feeding material into commercial models. Institutions that pursue this direction convert their integrity policy from a list of prohibitions into an expression of professional ethics, which the compliance literature has long suggested is the more durable foundation (Bertram Gallant, 2017).

6.3 Provenance, Watermarking, and Technical Standards

On the technical horizon, effort is shifting from post-hoc detection, which Section 4 argued is losing, to provenance: establishing at creation time where content came from. Statistical watermarking of model output has been demonstrated at scale (Kirchenbauer et al., 2023), and content-credential standards that cryptographically bind provenance metadata to documents are maturing in adjacent industries. The honest assessment is that these technologies will help but not solve: watermarks can be weakened by paraphrasing, they require cooperation from model providers, and open-weight models outside any provider's control will remain unwatermarked (Sadasivan et al., 2023; Lancaster, 2023). Their realistic role is to raise the cost and lower the deniability of wholesale outsourcing, while process evidence, version histories, drafting records, and interaction logs, supplies the corroboration that individual misconduct cases require. Institutions should track these standards and press vendors to support them, while resisting the temptation to treat any of them as the awaited silver bullet.

6.4 Implications for Indian and South Asian Higher Education

The stakes of getting these directions right are especially high in India and the wider South Asian region, home to one of the world's largest and fastest-growing higher-education systems. Several structural features amplify both the risks and the opportunities identified above. Enrolment scale means that assessment is dominated by high-volume written examinations and assignments, exactly the formats most exposed to GenAI completion, while student-to-faculty ratios leave little slack for the labour-intensive alternatives, such as vivas and process-based assessment, that smaller systems can adopt more readily. The National Education Policy's push toward continuous, competency-based evaluation creates a genuine policy window: institutions redesigning assessment anyway can build GenAI-resilience in from the start rather than retrofitting it later. At the same time, the equity concerns of Section 4.3 bite harder where connectivity, device



access, and paid-tier subscriptions are unevenly distributed across urban and rural cohorts, and where a large share of students write academic English as an additional language and therefore face elevated detector-bias risk (Liang et al., 2023; UNESCO, 2023).

Regulatory developments are beginning to respond. The University Grants Commission's academic integrity regulations, framed originally around plagiarism percentages in research submissions, illustrate the general problem this paper has described: similarity-based rules are silent on generated text, and mechanical thresholds transfer poorly to a world without source documents. Indian institutions consequently have an opportunity to leapfrog rather than repeat the detection-first missteps documented in Western systems, moving directly to disclosure-based policy, AI-literacy integration, and the layered governance model of Section 5. The region's substantial IT-audit and assurance profession, much of it already serving global clients on AI governance engagements, is a natural and largely untapped partner in that effort (Dwivedi et al., 2023; Crawford et al., 2023).

6.5 Research Agenda

The research literature, though expanding at remarkable speed, remains dominated by commentary, single-institution surveys, and benchmarking exercises (Lo, 2023; Bearman et al., 2023). Several gaps stand out. There is little longitudinal evidence on how sustained GenAI use affects learning itself, the question that ultimately underlies the integrity debate. There are few controlled evaluations of redesigned assessments' resistance to GenAI completion, and almost none of the fairness properties of integrity procedures in practice, including accusation rates across demographic groups. The governance dimension is nearly untouched: we lack studies of how institutions actually allocate ownership of GenAI risk, and audit-style evaluations of control effectiveness of the kind Figure 3 envisages. Cross-national work is also needed, since most published evidence comes from a handful of Anglophone systems while the majority of the world's students, including the rapidly digitising higher-education systems of South Asia, face the same pressures with different resource constraints (UNESCO, 2023; Zawacki-Richter et al., 2019).

VI. CONCLUSION

Generative AI has not destroyed academic integrity, but it has destroyed the enforcement model through which integrity was operationalised for a generation. Text-matching is blind to text that never existed before; AI detection is unreliable, evadable, and biased against the very students institutions owe the greatest care; and blanket prohibition is unenforceable in any unsupervised setting. Continuing to escalate along that path purchases the appearance of control at the price of fairness, trust, and credibility.

The alternative defended in this paper is neither surrender nor nostalgia. It treats integrity as an institutional risk domain governed through layered controls: principled policy with clear ownership, assessment and curriculum redesign that shrinks the attack surface while teaching the competence students will need anyway, and technology deployed as corroboration and provenance rather than as accuser, all bound together by a continuous cycle of monitoring, fair investigation, and independent review. For information governance and audit professionals, higher education's predicament is a preview of a question every organisation certifying human performance will soon face: how do we assure the authenticity of work in a world where machines write? Universities are answering it first, in public, with their credibility at stake. Getting the answer right, honesty about process, evidence over suspicion, and governance over gadgetry, matters far beyond the campus.

REFERENCES

1. Goyal, S., Rallabandi, B., Gattupalli, K. K., & Medarametla, M. S. (2025). Digital twin-enabled context-aware networks: Enhancing smart grid resilience through predictive intelligence. In 2025 2nd International Conference on Recent Trends in Electrical, Electronics and Computing Technologies (ICRTEECT) (pp. 1–6). IEEE. <https://doi.org/10.1109/ICRTEECT67512.2025.11448880>



2. Tripathi, N., Gattupalli, K. K., Rallabandi, B., & Medarametla, M. S. (2025). AI-native Cloud-RAN orchestration for enterprise private 5G using digital twin models. In 2025 1st IEEE Uttar Pradesh Section Women in Engineering International Conference on Electrical Electronics and Computer Engineering (UPWIECON) (pp. 851–857). IEEE. <https://doi.org/10.1109/UPWIECON67212.2025.11390348>
3. Goyal, S., Gattupalli, K. K., Rallabandi, B., & Medarametla, M. S. (2025). Integrating terrestrial and non-terrestrial networks for reliable rural coverage in agriculture and smart grids. In 2025 1st IEEE Uttar Pradesh Section Women in Engineering International Conference on Electrical Electronics and Computer Engineering (UPWIECON) (pp. 846–850). IEEE. <https://doi.org/10.1109/UPWIECON67212.2025.11390331>
4. N. Pachauri, V. Suresh, M. V. V. Prasad Kantipudi, R. Alkanhel, and H. A. Abdallah, “Multi-Drug Scheduling for Chemotherapy Using Fractional Order Internal Model Controller,” *Mathematics*, vol. 11, no. 8, Art. no. 1779, 2023,
5. M. V. V. Prasad Kantipudi, S. Vemuri, N. S. Pradeep Kumar, S. Sreenath Kashyap, and S. Eslamian, “Proposing Model for Water Quality Analysis Based on Hyperspectral Remote Sensor Data,” in *Handbook of Hydroinformatics*, Elsevier, 2023, pp. 317–324.
6. Rana, M., Srinivas, S., Jamili, L. K., Jaiswal, I. A., Nakka, S., & Kasetti, S. (2025, May). Real-Time Monitoring and Prediction of Blood Sugar Levels in Diabetic Patients with Functional Models. In 2025 International Conference on Engineering, Technology & Management (ICETM) (pp. 1-6). IEEE.
7. Jaiswal, I. A. (2023). Intelligent Cybersecurity Framework for Large-Scale RESTful Service Architectures. *International Journal of Research Radicals in Multidisciplinary Fields*, ISSN: 2960-043X, 2 (1), 178–184. Retrieved from <https://www.researchradicals.com/index.php/rr/article/view/252>.
8. Jaiswal, I. A. (2023). Multilingual and culturally adaptive AI models for global education platforms. *International Journal for Research in Education (IJRE)*, 12(9), 17–27. <https://doi.org/10.63345/ijre.v12.i9.1>
9. M. S. Prashanth, R. Aluvalu, and M. V. V. Kantipudi, “Enhancing Health Product Traceability on the Blockchain: A Novel Approach for Supply Chain Management Inspection to AI,” *EAI Endorsed Transactions on Pervasive Health & Technology*, vol. 10, no. 1, 2024.
10. H. K. Jani, M. V. V. Prasad Kantipudi, G. Nagababu, D. Prajapati, and S. S. Kachhwaha, “Simultaneity of Wind and Solar Energy: A Spatio-Temporal Analysis to Delineate the Plausible Regions to Harness,” *Sustainable Energy Technologies and Assessments*, vol. 53, Art. no. 102665, 2022
11. Jaiswal, I. A. (2022). Natural language processing for security policy and log analysis. *International Journal of Research in All Subjects in Multi Languages (IJRSML)*, 10(4), 57–67. <https://doi.org/10.63345/ijrsml.v10.i4.1>
12. Khan, S. (2019). Sales force security compliance: An in-depth study of GDPR, HIPAA, and PCI-DSS enforcement in cloud-based CRM systems and their implications for global enterprises. *International Journal of Advanced Research in Electrical, Electronics and Instrumentation Engineering*, 8(9), 2211–2219. <https://doi.org/10.15662/IJAREEIE.2019.0809010>
13. S. Choudhary, C. H. Kandikattu, S. Kumar, M. V. V. P. Kantipudi, and M. Kumar, “Enhancing cybersecurity through combined convolutional neural network-gated recurrent unit approach for distributed denial of service attack detection,” in *Proceedings of the 2024 1st International Conference on Innovative Engineering Sciences and Technological Research (ICIESTR)*, pp. 1–6, 2024.
14. S. Rani, D. Ghai, and S. Kumar, “Reconstruction of wire frame model of complex images using syntactic pattern recognition,” in *IET Conference Proceedings CP791*, vol. 2021, no. 11, pp. 8–13, 2021.
15. Swathi, S. Kumar, S. Rani, A. Jain, and R. K. M. V. N. M., “Emotion classification using feature extraction of facial expression,” in *Proceedings of the 2022 2nd International Conference on Technological Advancements in Computational Sciences (ICTACS)*, pp. 283–288, 2022.
16. S. Choudhary, S. Kumar, S. Gowroju, M. Gulhane, and R. Sri Lakshmi, Eds., *Genomics at the Nexus of AI, Computer Vision, and Machine Learning*. Hoboken, NJ, USA: John Wiley & Sons, 2024.



17. S. Kumar, S. Choudhary, S. Gowroju, and A. Bhola, "Convolutional neural network approach for multimodal biometric recognition system for banking sector on fusion of face and finger," in *Multimodal Biometric and Machine Learning Technologies: Applications for Computer Vision*, pp. 251–267, 2023.
18. S. Choudhary, S. Kumar, M. Gulhane, and M. Kumar, "Secured automated certificate creation based on multimodal biometric verification," in *Multimodal Biometric and Machine Learning Technologies: Applications for Computer Vision*, pp. 269–281, 2023.
19. Jaiswal, I. A. (2024). AI-powered observability and incident prediction in distributed enterprise platforms. *Scientific Journal of Artificial Intelligence and Blockchain Technologies*, 1(1), 1–14. <https://doi.org/10.63345/sjaibt.v1.i1.201>
20. S. Choudhary, C. H. Kandikattu, S. Kumar, M. V. V. P. Kantipudi, and M. Kumar, "Enhancing cybersecurity through combined convolutional neural network-gated recurrent unit approach for distributed denial of service attack detection," in *Proceedings of the 2024 1st International Conference on Innovative Engineering Sciences and Technological Research (ICIESTR)*, pp. 1–6, 2024.
21. Jaiswal, I. A. (2023). High-Performance AI-Augmented Content Management Systems for Distributed Clouds. *International Journal of Multidisciplinary Innovation and Research Methodology*, ISSN: 2960-2068, 2 (2), 90–97. Retrieved from <https://ijmirm.com/index.php/ijmirm/article/view/243>.
22. Khan, S. (2018). The role of zero-trust models in database security: Eliminating implicit trust and enforcing continuous verification in enterprise data access systems. *International Journal of Research in Electronics and Computer Engineering*, 6(1), 1622–1628.
23. Jaiswal, I. A. (2024). Self-healing REST services using artificial intelligence in multi-cloud environments. *Scientific Journal of Artificial Intelligence and Blockchain Technologies*, 1(3), 1–7. <https://doi.org/10.63345/sjaibt.v1.i3.201>
24. V. Saini, A. Jain, Anurag, and M. V. V. Prasad Kantipudi, "An Efficient Approach for Improving the Performance of Autonomous Vehicle Using Advanced Computer Vision," in *International Conference on Soft Computing and Pattern Recognition*, Cham, Switzerland: Springer Nature Switzerland, 2023, pp. 164–174.
25. Khan, S. (2017). The role of privacy-preserving techniques in database security: Secure multi-party computation, differential privacy, and homomorphic encryption. *The Research Journal*, 3(4), 29–35.
26. Khan, S. (2016). A study of data provenance and integrity in database security: Ensuring authenticity, non-repudiation, and accountability in data lifecycle management. *International Journal of Research in Electronics and Computer Engineering*, 4(3), 180–186.
27. Cherladine, K., Rallabandi, B. R., Goel, A., Kaushik, K., Dhillon, A., & Soni, M. (2025). Generative adversarial defense models for simulated cyberattack scenarios in virtual networks. In *2025 International Conference on Digital Innovations for Sustainable Solutions (ICDISS)* (pp. 1–6). IEEE. <https://doi.org/10.1109/ICDISS68238.2025.11320686>
28. Rallabandi, B. R. (2020). MEC-native 5G systems: Orchestration algorithms for ultra-low latency cloud-edge integration. *International Journal of Intelligent Systems and Applications in Engineering*, 8(4), 398–408. <https://ijisae.org/index.php/IJISAE/article/view/7889>
29. Rallabandi, B. R. (2018). Joint deployment and operational energy optimization in heterogeneous cellular networks under traffic variability. *International Journal of Communication Networks and Information Security*, 10(3), 638–647. <https://ijcnis.org/index.php/ijcnis/article/view/8604>
30. Kalal, M. (2026, February). Predictive Production Planning in Advanced Manufacturing Using SAP PP and Analytics. In *2026 5th International Conference on Sentiment Analysis and Deep Learning (ICSADL)* (pp. 1363-1370). IEEE.
31. M. Kalal, "Integrating SAP PP and SAP Digital Manufacturing for Lean Production Optimization," 2026 International Symposium of Systems, Advanced Technologies and Knowledge (ISSATK), Hammamet, Tunisia, 2026, pp. 1-7, doi: 10.1109/ISSATK69667.2026.11566800.



32. M. Kalal, Y. R. Krishna, B. Jayswal, B. Konduru and V. S. A. Kondru, "Metaheuristic Optimization-Driven Information Management System for Enterprise Intelligence," 2026 IEEE 15th International Conference on Communication Systems and Network Technologies (CSNT), Al-Khobar, Saudi Arabia, 2026, pp. 1417-1423, doi: 10.1109/CSNT69054.2026.11502494.
33. M. Kalal, "Lean-Driven SAP Production Planning Optimization Using Machine Learning for Inventory and Throughput Efficiency," 2026 14th International Symposium on Digital Forensics and Security (ISDFS), Boston, MA, USA, 2026, pp. 1-6, doi: 10.1109/ISDFS69419.2026.11459029.
34. Kalal, M. (2024). Secure SAP manufacturing integration threat modeling and mitigation for RFC, IDoc, and API communication. International Journal of Communication Networks and Information Security, 16(4). <https://doi.org/10.5281/zenodo.20088018>
Kalal, M., & Tanner, O. C. (2025). Autonomous exception management in SAP S/4HANA manufacturing through multi-agent generative AI and event-driven supply networks. International Journal of Core Engineering & Management, 8(4), 130–137.
35. C. H. Neetha and M. P. Kantipudi, "Adaptive Edge-Based Bilinear Interpolation for Smart Healthcare," in IET Conference Proceedings CP824, vol. 2022, no. 26, Stevenage, U.K.: The Institution of Engineering and Technology (IET), 2022, pp. 7–13.
36. R. Aluvalu, R. Venkatachalam, V. Uma Maheswari, R. J. Anandhi, M. V. V. Kantipudi, and S. Prashanth Mallellu, "IWHO-Based Cluster Head Selection for Vehicle-to-Vehicle Communication in Intelligent Transport System," International Journal of Transport Development and Integration, vol. 8, no. 2, pp. 154–166, 2024
37. S. Velamuri, S. K. Sudabattula, M. V. V. Prasad Kantipudi, and N. Prabakaran, "Q-Learning Based Commercial Electric Vehicles Scheduling in a Renewable Energy Dominant Distribution Systems," Electric Power Components and Systems, pp. 1–14, 2023 Singh, M., Rallabandi, B., Gattupalli, K. K., & Sai Medarametla, M. (2025). Autonomous private 5G field area networks: Achieving secure, scalable, and self-healing smart grid communications. In 2025 2nd International Conference on Recent Trends in Electrical, Electronics and Computing Technologies (ICRTEECT) (pp. 1–6). IEEE. <https://doi.org/10.1109/ICRTEECT67512.2025.11448712>
38. Desai, A. B., Rallabandi, B., Gattupalli, K. K., & Medarametla, M. S. (2025). Agentic AI for rural connectivity and spectrum-aware networks to power smart agriculture ecosystems. In 2025 2nd International Conference on Recent Trends in Electrical, Electronics and Computing Technologies (ICRTEECT) (pp. 1–6). IEEE. <https://doi.org/10.1109/ICRTEECT67512.2025.11448879>

