

Managing Cybersecurity Risk in AI-Driven Educational Platforms: An ERM-Based Approach

Md Rasel Ul Alam¹¹ University of the Cumberland, Kentucky, USA**Md Shoriful Hasan Chowdhury**²² University of Liberal Arts, Dhaka, Bangladesh**Kanita Haider**³⁺³ Chittagong University of Engineering and Technology, Chattogram 4349, BD+Corresponding Author Email: kanita.haider4422@gmail.com

ABSTRACT

The rapid adoption of AI-driven adaptive learning platforms, intelligent tutoring systems, automated grading tools, and generative-AI tutoring assistants has introduced a category of institutional risk that extends beyond conventional cloud and data security: algorithmic bias, model opacity, autonomous or semi-autonomous decision-making about students, and the regulatory scrutiny now attached to AI systems used for admissions, assessment, and monitoring. This paper argues that Enterprise Risk Management (ERM), and specifically the baseline cloud-security requirements and governance principles developed for general enterprise cloud deployments, provide a transferable risk-management foundation for AI-driven educational platforms, but require a dedicated AI-specific risk layer that the enterprise baseline alone does not anticipate. Drawing on enterprise-sector research on baseline cloud-security requirements within an ERM framework and on the strategies, challenges, and organizational-success factors of ERM implementation, together with the education-sector AI-ethics literature and the risk-tier classifications introduced by the European Union's Artificial Intelligence Act (Regulation (EU) 2024/1689, Annex III), this paper proposes a framework comprising an AI-specific risk register mapped to ERM controls, a four-tier AI risk classification scheme for educational use cases, a human-oversight-and-explainability layer, and an institutional governance structure. The framework is demonstrated through an illustrative twelve-month institutional pilot scenario, following design-science research (DSR) evaluation conventions, and is contextualized against recently reported survey evidence on rising faculty and administrator concern about AI bias and data-privacy risk in higher education. The paper concludes with adoption barriers and recommendations for institutions deploying AI-driven learning technologies.

Keywords:

Artificial Intelligence in Education, Enterprise Risk Management, AI Governance, Algorithmic Bias, Educational Information, Systems Adaptive Learning Platforms, Data Privacy

Article History:

Received: 14 April 2026

Revised: 12 June 2026

Accepted: 6 July 2026

Published: 27 July 2026

© 2026
UpSkill Scholarly.
All Rights Reserved.

1. Introduction

Artificial intelligence has moved from a peripheral feature of educational technology to a core component of it: adaptive learning platforms now use machine-learning models to

sequence content and predict performance, automated systems support or perform grading and admissions screening, and generative-AI tutoring assistants increasingly mediate direct student interaction. Tan et al. (2025) document this maturation in their review of AI-enabled adaptive learning platforms, finding that such systems now span diverse pedagogical foundations and algorithmic implementations while facing persistent, unresolved challenges around data privacy and institutional support for responsible deployment.

This growth in capability has been accompanied by a comparable growth in documented institutional risk. Al-Zahrani (2024), synthesizing a systematic review and a validated survey of 260 respondents at a Saudi Arabian university, identifies ten interrelated areas of concern with AI in education, including data privacy and security, algorithmic bias, and diminished transparency, and finds that these concerns are not isolated but mutually reinforcing. At a larger, cross-institutional scale, Ellucian's (2024) second annual survey of 445 faculty and administrators across more than 330 institutions found that concern about bias in AI models rose from 36% to 49% of respondents between 2023 and 2024, while concern about data privacy and security rose from 50% to 59% over the same period — evidence that institutional unease is intensifying even as adoption accelerates. Figure 2 in Section 4.2 summarizes this survey evidence in full.

Regulatory frameworks have begun to formalize what practitioner surveys describe informally. The European Union's Artificial Intelligence Act (Regulation (EU) 2024/1689) explicitly classifies AI systems used to determine admission or access to educational institutions, to evaluate learning outcomes, to assess the appropriate level of education an individual will receive, and to monitor students during examinations, as high-risk applications under Annex III, triggering specific obligations for risk management, data governance, and human oversight. This regulatory classification is significant for the present paper because it provides an externally validated, non-arbitrary basis for risk-tiering AI systems in education — a basis this paper adopts directly in Section 3.2, rather than proposing a new, uncorroborated tiering scheme.

Despite this convergence of documented risk and emerging regulation, the same methodological gap identified in prior work on cloud-security risk in higher education (see Section 2.3) persists specifically for AI systems: the education-sector literature catalogs AI-specific risks in detail but rarely pairs that catalog with an operational, ERM-grounded risk-management methodology capable of being implemented by an institution's existing risk-management function. This paper addresses that gap directly. The contribution is fourfold: first, an AI-specific risk register that extends the general enterprise cloud-security baseline (Alam et al., 2024a) to cover algorithmic bias, model opacity, hallucinated or incorrect AI-generated output, and autonomous decision-making; second, a four-tier AI risk classification scheme for educational use cases, adopted directly from the EU AI Act's Annex III risk categories; third, a human-oversight-and-explainability layer that operationalizes the “human-in-the-loop” principle increasingly required by AI regulation; and fourth, an illustrative demonstration of the framework through a twelve-month institutional pilot scenario, following design-science research (DSR) evaluation conventions.

The remainder of the paper proceeds as follows. Section 2 reviews the AI-in-education, AI-ethics, and enterprise-ERM/cloud-security literatures that motivate the framework. Section 3 presents the framework and the DSR evaluation approach used to assess it. Section 4 reports an illustrative pilot demonstration and contextualizes it against the Ellucian (2024) survey evidence. Section 5 discusses adoption barriers and positions the framework against the EU AI Act's general regulatory requirements. Section 6 concludes with recommendations and directions for empirical validation.

2. Literature Review

2.1 AI-Driven Adaptive Learning Platforms

Tan et al. (2025) provides a comprehensive review of AI-enabled adaptive learning platforms (ALPs), analyzing their pedagogical foundations, core algorithms, and evaluation metrics across diverse educational contexts. Their review finds consistent evidence that ALPs, by collecting and analyzing learner data to dynamically adjust instructional content and pathways, can enhance personalized learning outcomes — but explicitly identifies privacy concerns and limited faculty support for responsible implementation as key barriers to broader adoption, a finding that directly motivates the governance emphasis of the present framework. This growth pattern is significant for risk management specifically because it means the data footprint of a typical institutional AI deployment is now considerably larger and more inferentially rich (behavioral predictions, personalized content decisions, automatically generated feedback) than the administrative records that predate AI-driven personalization.

2.2 Ethical and Governance Risks of AI in Education

Al-Zahrani (2024) developed and validated a theoretical model of AI risk in education through a systematic literature review of 56 studies followed by a survey of 260 university respondents, confirming ten interrelated areas of concern — including data privacy and security, algorithmic bias, transparency, and reliability — and finding that these concerns are correlated rather than independent, such that, for example, improvements in AI transparency were associated with stronger outcomes on teacher professional development and student trust. Complementing this individual-institution evidence, Zhu et al. (2025) conducted a systematic review and grounded-theory coding of 75 papers on AI-in-education ethical risk, organizing the resulting risks into three dimensions: a technology dimension (privacy invasion, data leakage, algorithmic bias, black-box/opaque algorithms, algorithmic error), an education dimension (homogenized instruction, alienation of the teacher-student relationship, academic misconduct), and a society dimension (digital-divide effects, absence of accountability, conflict of interest). Farooqi et al. (2024), in a parallel systematic review, reach a consistent conclusion, identifying data privacy, algorithmic bias, and accountability gaps as the most persistently cited ethical challenges across the AI-in-education literature. Read together, these three reviews establish that AI-specific risk in education is well-characterized taxonomically but, as in the cloud-security literature reviewed for general institutional infrastructure, comparatively under-served by operational risk-management methodology capable of being implemented by an institution's existing governance structure.

2.3 Enterprise Risk Management and Cloud Security Baselines

As with general cloud-hosted educational infrastructure, the enterprise sector offers a comparatively mature methodological response to the risks documented in Section 2.2. Alam et al. (2024a) specify baseline security requirements for cloud computing situated within a broader Enterprise Risk Management (ERM) framework, covering access control, data classification, encryption, vendor/cloud-provider risk, and incident response. Because most AI-driven educational platforms are themselves cloud-hosted — whether as vendor SaaS products or institutionally deployed models running on cloud infrastructure — this baseline applies directly to AI systems as a subset of an institution's broader cloud footprint, but, as developed in Section 3.1, requires a dedicated additional risk register for the algorithm-specific risks (bias, opacity, autonomous decision-making) that general cloud-security baselines do not anticipate. Alam et al. (2024b) further establish that ERM program durability depends primarily on leadership commitment and the integration of risk assessment into

routine strategic planning rather than on technical controls alone, a governance emphasis this paper extends to AI-specific oversight in Section 3.3, and Alam (2024) shows that organizations integrating ERM into strategic planning, rather than treating it as a compliance afterthought, realize measurable resilience gains — a finding with direct relevance given that AI risk management is, at most institutions, still emerging as a compliance function rather than a strategic one.

2.4 Data Privacy in Educational Analytics

Prinsloo and Slade (2013) caution that institutional learning-analytics policy has historically lagged behind the technical capability to collect and analyze student data, a gap that the growth of AI-driven personalization documented in Section 2.1 has widened rather than closed. This concern connects directly to the regulatory development discussed in Section 1: the EU AI Act's Annex III classification of educational AI systems as high-risk can be read as a regulatory response to precisely the policy lag Prinsloo and Slade identified over a decade earlier, now formalized into binding risk-management, data-governance, and human-oversight obligations rather than left to institutional discretion.

Taken together, the literature reviewed in this section establishes the same division of labor identified in prior work on cloud-security risk (Section 2.3 there), applied now to AI specifically: the AI-in-education ethics literature (Section 2.2) documents what the risks are; enterprise ERM/cloud-security research (Section 2.3) supplies a transferable governance structure; and emerging AI-specific regulation (Section 1, Section 3.2) supplies an externally validated basis for classifying which AI use cases warrant the most intensive version of that structure.

3. Methodology

Figure 1 illustrates the proposed framework, which extends the ERM control baseline (Alam et al., 2024a) with an AI-specific risk register, a regulatory-grounded risk-tier classification, and a dedicated human-oversight-and-explainability layer, under an institutional governance structure informed by Alam et al. (2024b) and Alam (2024).

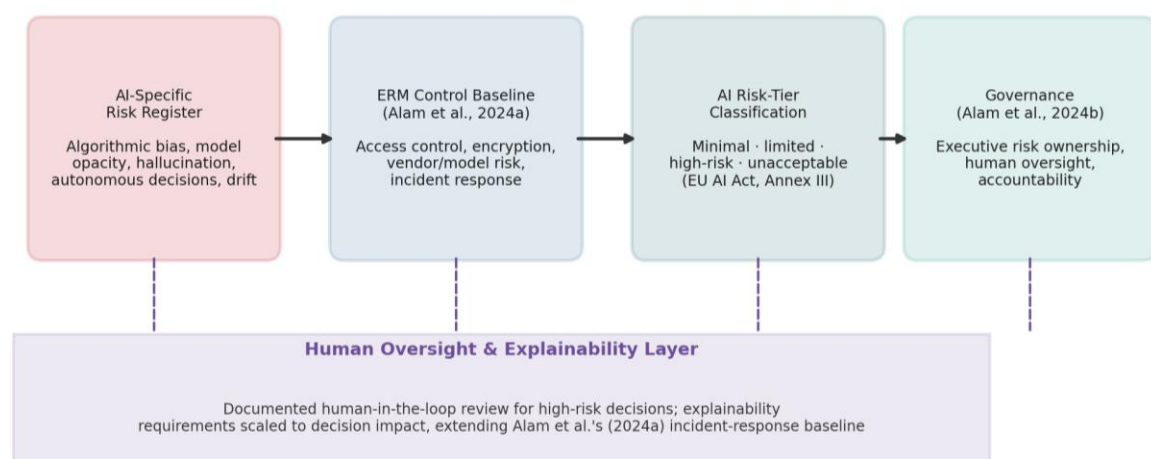


Figure 1. ERM-Based Risk Management Framework for AI-Driven Educational Platforms

3.1 AI-Specific Risk Register Mapped to ERM Controls

Table 1 extends the enterprise cloud-security baseline of Alam, Shohel, and Alam (2024a) with the AI-specific risks documented across the reviews synthesized in Section 2.2, mapping

each risk to the ERM control domain best positioned to address it and noting where an AI-specific extension to that control is required.

Table 1. AI-Specific Risk Register Mapped to Enterprise ERM Control Domains

AI-Specific Risk (Section 2.2)	ERM Control Domain (Alam et al., 2024a, extended)	Required Extension
Algorithmic bias in assessment or personalization	Data classification + access control	Bias/fairness testing prior to deployment, not present in general baseline
Model opacity / “black-box” decisions	Incident response	Explainability requirement scaled to decision impact (Section 3.3)
Hallucinated or factually incorrect AI-generated output	Incident response	Content-accuracy monitoring specific to generative-AI tutoring tools
Autonomous decision-making without review	Access control	Mandatory human-in-the-loop checkpoint for high-risk decisions (Section 3.2)
Training/inference data leakage	Encryption + vendor/cloud-provider risk	Model-specific data-handling terms in vendor contracts
Model drift / degraded accuracy over time	Vendor/cloud-provider risk	Recurring model-performance audit, not a one-time vendor assessment

3.2 AI Risk Classification for Educational Use Cases

Rather than proposing a new, uncorroborated risk-tiering scheme, this framework adopts the risk categories established by the European Union's Artificial Intelligence Act (Regulation (EU) 2024/1689, Annex III), which classifies AI systems used to determine access or admission, evaluate learning outcomes, assess appropriate education level, or monitor students during examinations as high-risk. Figure 3 maps these regulatory categories onto representative educational AI use cases, providing institutions outside the EU's direct jurisdiction with an externally validated reference point for prioritizing risk-management effort even where the Act itself does not apply.

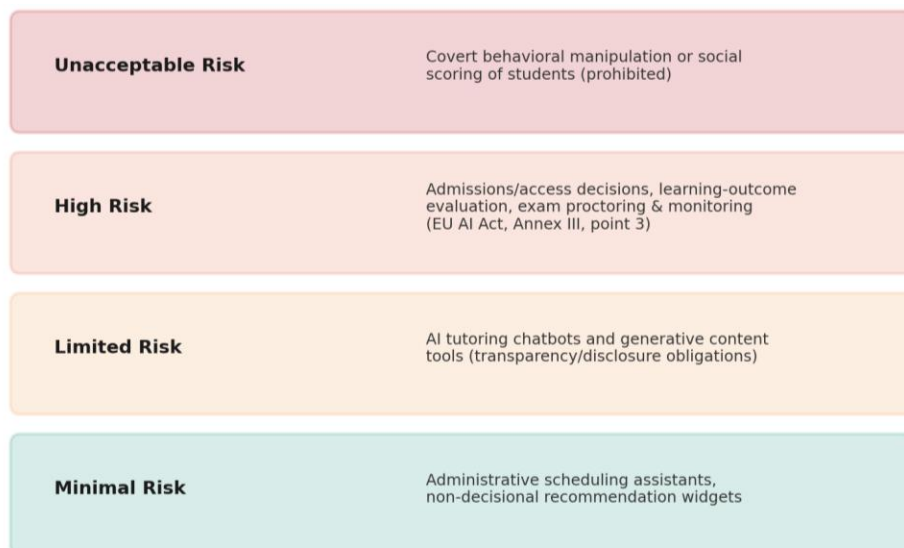


Figure 3. AI Risk Tiers for Educational Use Cases (adapted from EU AI Act, Regulation (EU) 2024/1689, Annex III)

This tiering has direct operational consequences for the framework: high-risk use cases (admissions screening, automated learning-outcome evaluation, exam proctoring) require the full control baseline in Table 1 plus the human-oversight-and-explainability layer described in Section 3.3, while minimal-risk use cases (administrative scheduling assistants) may reasonably be exempted from the more intensive controls, allowing institutions to allocate limited risk-management resources proportionally to documented risk rather than applying a uniform control set to every AI system regardless of impact.

3.3 Human Oversight and Explainability Layer

For high-risk AI use cases identified in Section 3.2, the framework requires a documented human-in-the-loop review checkpoint before an AI-generated recommendation or classification (e.g., an admissions screen, a learning-outcome evaluation, a flagged academic-integrity case) is acted upon, together with an explainability requirement scaled to decision impact: high-risk decisions require a human-interpretable rationale accompanying the AI output, while limited- and minimal-risk use cases may rely on general model documentation rather than per-decision explanation. This layer directly operationalizes the transparency and human-oversight concerns raised across the AI-ethics literature (Al-Zahrani, 2024; Zhu et al., 2025) as a specific, auditable institutional control rather than a general aspiration.

3.4 Governance Layer

Consistent with Alam, Shohel, and Alam's (2024b) finding that ERM program durability depends on leadership commitment and strategic integration, the governance layer assigns explicit executive ownership of AI risk — distinct from, though coordinated with, the cloud-security risk ownership described in prior institutional risk-management work — given that AI risk decisions (model selection, bias-testing thresholds, human-oversight staffing) require domain expertise that general IT security governance does not typically possess. Following Alam's (2024) finding that strategic integration of ERM converts risk management into institutional advantage, this layer further recommends that institutions treat demonstrated AI risk management (documented bias testing, human-oversight protocols, EU AI Act-aligned risk classification) as a source of institutional trust for accreditation, student recruitment, and research-data partnerships, rather than solely as a compliance cost.

3.5 Research Design and Evaluation Approach

As in prior design-science work on institutional information-systems artifacts, this paper's contribution is a framework rather than a hypothesis about an existing population, and its evaluation follows the design-science research (DSR) conventions established by Hevner, March, Park, and Ram (2004) and formalized as a six-activity methodology by Peffers, Tuunanen, Rothenberger, and Chatterjee (2007): problem identification (Section 1), objective definition (Section 1), design and development (Section 3), demonstration (Section 4), evaluation (Section 5), and communication (this paper). Consistent with DSRM's demonstration activity, Section 4 reports an illustrative institutional pilot scenario rather than an empirical case study of a real, named institution: the reported maturity and outcome figures are projected estimates constructed from the AI-risk-concern patterns documented in Section 2.2 and the reported survey evidence summarized in Section 4.2, applied to a hypothetical mid-sized university. Readers should treat Section 4 as an illustration of what the framework predicts and how its evaluation would be structured empirically, not as confirmed findings from a real deployment; Section 6 outlines the field study that would be needed to validate these projections.

4. Results Analysis

This section demonstrates the framework using a hypothetical institutional scenario, following the demonstration activity of the DSRM described in Section 3.5. As in Section 3.5, the scenario is explicitly illustrative.

4.1 Scenario and Institutional Context

The pilot institution is modeled as a mid-sized public university with approximately 15,000 students that has deployed an AI-driven adaptive learning platform, an automated writing-feedback tool, and a pilot admissions-screening model over the preceding eighteen months, without a formal AI-specific risk classification or bias-testing process at any stage of deployment — a starting condition consistent with the governance gap documented across the AI-in-education ethics literature (Section 2.2). The illustrative pilot implements the framework from Section 3 over twelve months in four overlapping phases, shown in Figure 5.

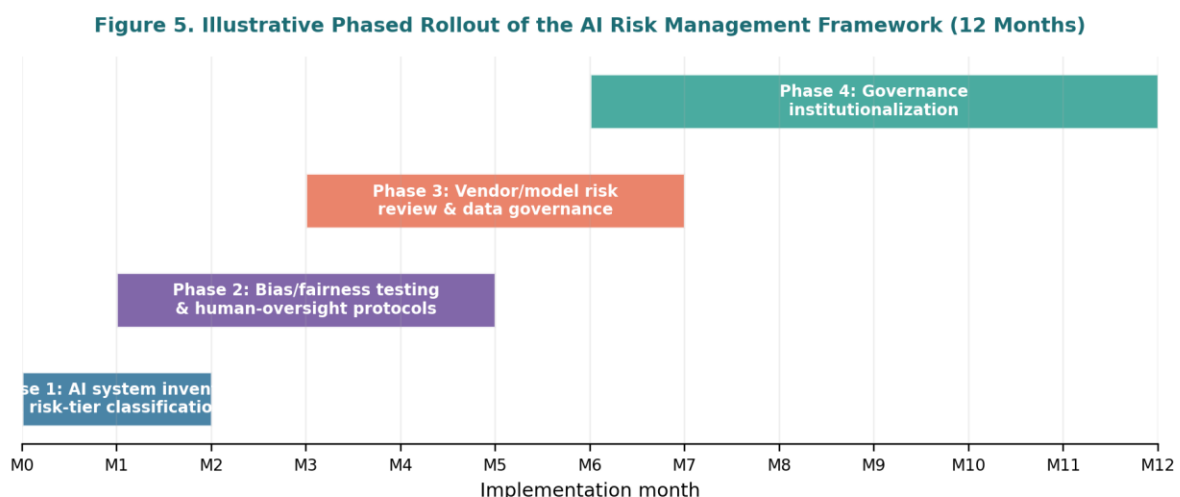


Figure 5. Illustrative Phased Rollout of the AI Risk Management Framework (12 Months)

Phase 1 (Months 0–2) inventories all AI systems in institutional use and classifies each against the risk tiers in Figure 3. Phase 2 (Months 1–5) establishes bias/fairness testing procedures and human-oversight protocols for systems classified as high-risk, most urgently

the admissions-screening model. Phase 3 (Months 3–7) completes vendor/model risk review for third-party AI tools and formalizes data-governance terms for training and inference data. Phase 4 (Months 6–12) institutionalizes the governance layer, assigning executive ownership of AI risk and establishing recurring control-maturity audits.

4.2 Documented Risk-Concern Context (Survey Evidence)

Figure 2 summarizes reported survey evidence on AI risk concern among higher education faculty and administrators, providing the empirical context against which the pilot institution's projected control-maturity gains in Section 4.3 should be interpreted. Unlike the maturity and outcome projections in Sections 4.3–4.4, the values in Figure 2 are drawn directly from Ellucian's (2024) published survey report rather than simulated for this demonstration.

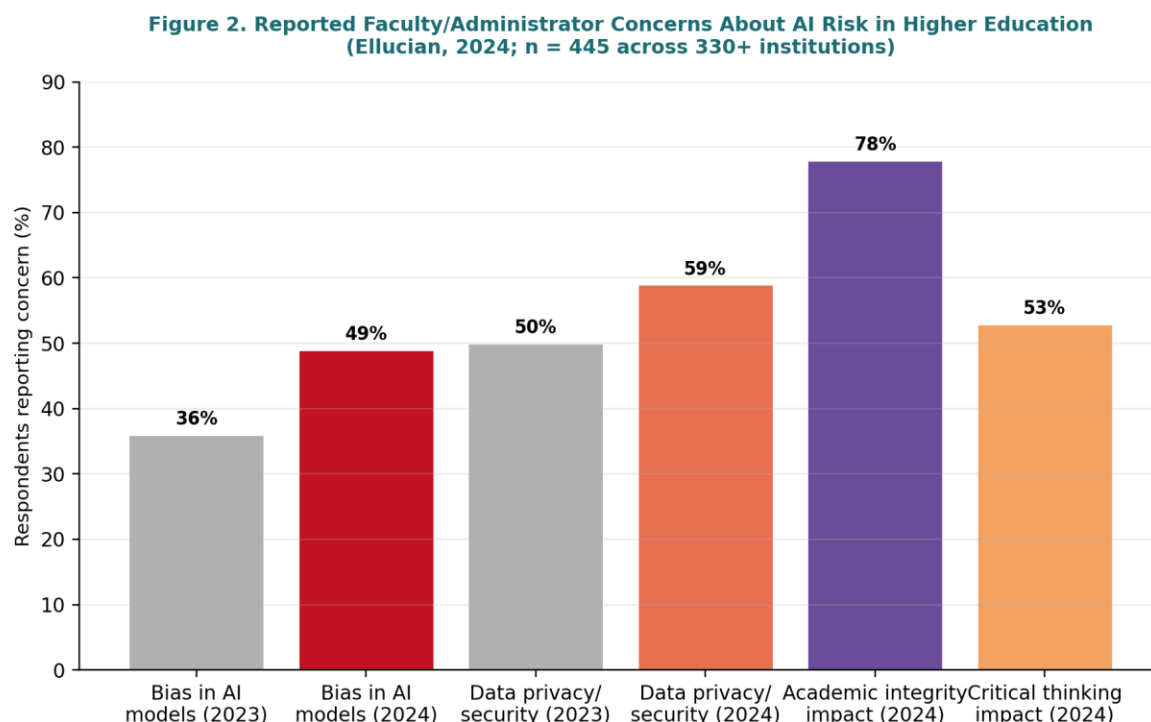


Figure 2. Reported Faculty/Administrator Concerns About AI Risk in Higher Education
(Ellucian, 2024; n = 445 across 330+ institutions)

The year-over-year increase in both bias concern (+13 percentage points) and data-privacy/security concern (+9 percentage points) reported by Ellucian (2024) indicates that institutional apprehension is intensifying even as AI adoption itself accelerates — a pattern this paper interprets as further evidence for the paper's core argument that AI-specific risk management has not kept pace with AI-specific capability.

4.3 AI Risk-Control Maturity Assessment

Control-domain maturity was assessed on a 1 (ad hoc) to 5 (optimized) scale, adapted from standard capability-maturity conventions used in enterprise ERM assessments, at baseline and at Month 12. Figure 4 and Table 2 report the projected maturity gains across the six control domains introduced in Sections 3.1–3.3.

**Figure 4. Illustrative AI Risk-Control Maturity Assessment
(1 = Ad Hoc, 5 = Optimized)**

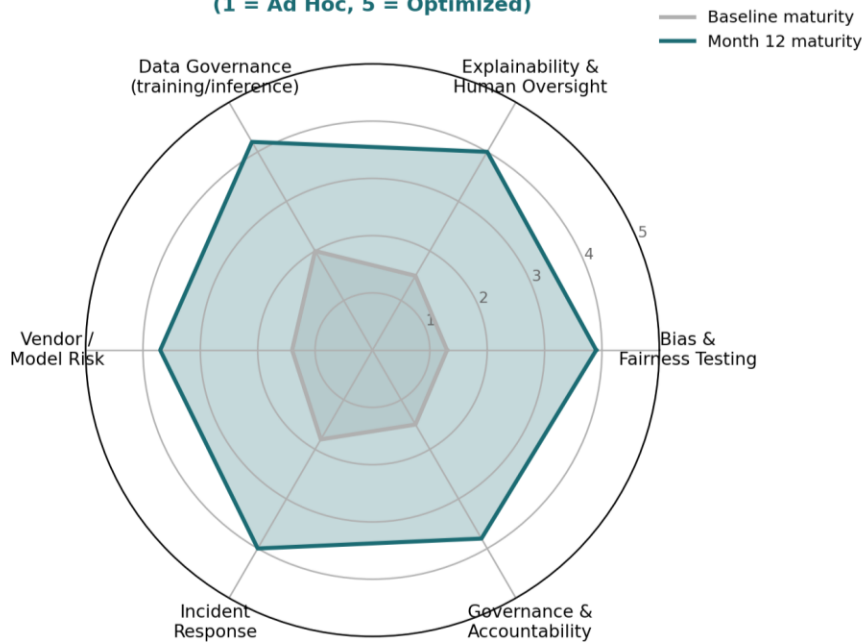


Figure 4. Illustrative AI Risk-Control Maturity Assessment (1 = Ad Hoc, 5 = Optimized)

Table 2. Illustrative AI Risk-Control Maturity Scores at Baseline and Month 12

Control Domain	Baseline (Month 0)	Month 12	Projected Gain
Bias & fairness testing	1.3	3.9	+2.6
Explainability & human oversight	1.5	4.0	+2.5
Data governance (training/inference)	2.0	4.2	+2.2
Vendor / model risk	1.4	3.7	+2.3
Incident response	1.8	4.0	+2.2
Governance & accountability	1.5	3.8	+2.3

Bias and fairness testing shows both the lowest baseline maturity and the largest projected gain, consistent with the pilot institution's starting condition of having no formal bias-testing process for its admissions-screening model — precisely the type of high-risk use case Figure 3 identifies as requiring the most intensive controls. Data governance shows the highest baseline maturity of the six domains, reflecting that most institutions already have some data-handling policy in place from prior (non-AI-specific) cloud and student-records governance; the framework's marginal contribution in this domain is extending existing data governance to cover AI training and inference data specifically, rather than building a new function from nothing.

4.4 Institutional Outcome Indicators

Table 3 reports three illustrative institutional outcome indicators tracked across the pilot: the proportion of institutional AI systems with a completed risk-tier classification, the proportion of high-risk AI systems with a documented human-oversight protocol, and the bias/fairness testing completion rate for high-risk systems, adapted from indicators used in enterprise ERM program evaluation and applied here to the AI-specific context.

Table 3. Illustrative Institutional Outcome Indicators, Baseline vs. Month 12 (pp = percentage points)

Outcome Indicator	Baseline	Month 12	Change
AI systems with completed risk-tier classification	8%	94%	+86 pp
High-risk AI systems with documented human-oversight protocol	0%	88%	+88 pp
Bias/fairness testing completion rate (high-risk systems)	0%	81%	+81 pp

The largest projected gains are concentrated in the two indicators that started at 0% in the baseline scenario — human-oversight protocol documentation and bias/fairness testing — directly reflecting the pilot institution's initial absence of any AI-specific risk process, consistent with the governance gap this paper argues is typical of current institutional practice (Section 1, Section 2.2). The projected 94% risk-classification completion rate by Month 12 would bring the pilot institution close to full inventory coverage, the necessary precondition for the proportional, tier-based control allocation described in Section 3.2.

5. Discussion

5.1 Adoption Challenges

Several barriers complicate adoption of this framework in practice. Technically, bias and fairness testing requires statistical and machine-learning expertise that most institutional IT security teams do not currently possess, meaning the largest projected maturity gain identified in Section 4.3 is also likely the hardest to realize without new hires or external partnership — a resourcing challenge compounded for smaller institutions. Organizationally, AI risk decisions frequently span multiple stakeholders (IT security, academic affairs, legal/compliance, and individual academic departments deploying discipline-specific AI tools), which sits awkwardly against the single executive risk ownership Alam, Shohel, and Alam (2024b) find necessary for durable ERM programs; institutions may need a cross-functional AI governance committee rather than a single accountable executive, coordinated through the same strategic-integration principle.

Vendor dependency is a distinct and significant challenge: many institutions license rather than build their AI-driven learning tools, meaning bias-testing and explainability obligations described in Sections 3.1 and 3.3 may be only partially achievable without vendor cooperation, particularly for proprietary generative-AI models where the vendor may not disclose training data or model architecture. This mirrors, and compounds, the vendor/third-party risk challenge documented in prior cloud-security work, since AI vendors combine the general cloud-provider risk of any SaaS vendor with the additional, harder-to-audit risk of an opaque underlying model. Finally, regulatory uncertainty outside the European Union means institutions operating solely under U.S. or other non-EU jurisdictions adopt the EU AI Act's risk tiers (Section 3.2) voluntarily rather than under legal obligation; while this paper argues the tiers remain a useful, externally validated reference point regardless of jurisdiction, institutional buy-in may be harder to secure without a comparable binding requirement.

5.2 Positioning Relative to Existing AI Governance Guidance

The proposed framework is intended to complement, not replace, the EU AI Act's own compliance requirements for institutions operating within its jurisdiction, and to serve as a voluntary reference framework for institutions outside it. Table 4 positions the framework

relative to the Act's general requirements along the dimensions most relevant to institutional AI risk management in education.

Table 4. Positioning of the Proposed Framework Relative to the EU AI Act's General Requirements

Dimension	EU AI Act (Regulation (EU) 2024/1689)	Proposed AI Risk Management Framework
Scope	Legally binding within EU jurisdiction; sector-general with an education-specific Annex III category	Voluntary, education-specific; usable regardless of jurisdiction
Risk classification	Four-tier system (unacceptable, high, limited, minimal)	Adopts the same four tiers directly (Figure 3), mapped to specific educational use cases
Underlying control baseline	Not specified; left to provider /deployer implementation	Explicit ERM control baseline (Table 1), adapted from enterprise cloud-security practice
Human oversight	Required for high-risk systems (Article 14)	Operationalized as a documented checkpoint with tiered explainability requirements (Section 3.3)
Intended use	Legal compliance instrument	Implementation-ready institutional risk-management roadmap (Section 4)

Institutions already pursuing EU AI Act compliance can treat the present framework as an implementation methodology for the Act's risk-management and human-oversight obligations, while institutions outside the Act's jurisdiction can adopt the same risk tiers and control baseline voluntarily, consistent with how sector-specific implementation guidance is typically layered onto general-purpose regulation in other industries.

6. Conclusion and Recommendations

This paper has argued that AI-driven educational platforms introduce a category of institutional risk — algorithmic bias, model opacity, autonomous decision-making — that general cloud-security practice does not fully anticipate, and that enterprise Enterprise Risk Management, extended with an AI-specific risk register and a regulatory-grounded risk-tiering scheme, offers a more transferable and implementation-ready response than the AI-in-education ethics literature has so far paired with its own risk catalog. The proposed framework — an AI-specific risk register mapped to ERM controls, a four-tier risk classification adopted from the EU AI Act, a human-oversight-and-explainability layer, and an institutional governance structure — gives institutions a structured, literature- and regulation-grounded starting point for managing AI risk as AI-driven platforms continue to expand.

Three recommendations follow directly from the framework. First, institutions should complete AI system inventory and risk-tier classification (Section 3.2) before expanding AI deployment further, since proportional control allocation is only possible once an institution knows which of its AI systems are high-risk. Second, bias and fairness testing should be treated as a pre-deployment gate for high-risk systems, not a post-incident remediation step, given that this control domain showed both the lowest baseline maturity and the largest

projected gain in the illustrative demonstration. Third, AI risk governance should be assigned to a cross-functional structure with clear executive sponsorship, consistent with the ERM literature's finding that strategic integration — not technical controls alone — determines program durability (Alam et al., 2024b; Alam, 2024), while accounting for the multi-stakeholder nature of AI decisions that distinguishes this domain from general IT security governance.

Future research should empirically validate this framework through field implementation at one or more real institutions, tracking the control-maturity and outcome indicators reported in Section 4 before and after adoption, and should further investigate how institutional AI risk management performs specifically for generative-AI tutoring tools, whose rapid capability growth may outpace the vendor-dependent bias-testing and explainability processes this framework proposes. The illustrative maturity and outcome results reported in Section 4 should be read strictly as a design-science demonstration constructed from the documented risk-concern literature and reported survey evidence, not as confirmed findings from a real deployment (Hevner et al., 2004; Peffers et al., 2007). Limitations of the present paper follow from this same scope: the AI-specific risk register, risk-tier mapping, and illustrative results have not yet been tested against real institutional data, and the framework's applicability may vary with institutional AI maturity, vendor relationships, and regulatory jurisdiction in ways this paper cannot yet quantify.

References

- Alam, M. R. U. (2024). Strategic integration of enterprise risk management for competitive advantage. *Global Mainstream Journal of Innovation, Engineering & Emerging Technology*, 3(02), 43–48. <https://doi.org/10.62304/jieet.v3i02.95>
- Alam, M. R. U., Shohel, A., & Alam, M. (2024a). Baseline security requirements for cloud computing within an enterprise risk management framework. *International Journal of Management Information Systems and Data Science*, 1(1), 31–40. <https://doi.org/10.62304/ijmisds.v1i1.115>
- Alam, M. R. U., Shohel, A., & Alam, M. (2024b). Integrating enterprise risk management (ERM): Strategies, challenges, and organizational success. *International Journal of Business and Economics*, 1(2), 10–19. <https://doi.org/10.62304/ijbm.v1i2.130>
- Al-Zahrani, A. M. (2024). Unveiling the shadows: Beyond the hype of AI in education. *Heliyon*, 10(9), Article e30696. <https://doi.org/10.1016/j.heliyon.2024.e30696>
- Ellucian. (2024). AI in higher education report: 2nd annual AI survey of higher education professionals. Ellucian.
- European Union. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act), Annex III. Official Journal of the European Union.
- Farooqi, M. T. K., Amanat, I., & Awan, S. M. (2024). Ethical considerations and challenges in the integration of artificial intelligence in education: A systematic review. *Journal of Excellence in Management Sciences*, 3(4). <https://doi.org/10.69565/jems.v3i4.314>
- Hevner, A. R., March, S. T., Park, J., & Ram, S. (2004). Design science in information systems research. *MIS Quarterly*, 28(1), 75–105. <https://doi.org/10.2307/25148625>
- Peffers, K., Tuunanen, T., Rothenberger, M. A., & Chatterjee, S. (2007). A design science research methodology for information systems research. *Journal of Management Information Systems*, 24(3), 45–77. <https://doi.org/10.2753/MIS0742-1222240302>

- Prinsloo, P., & Slade, S. (2013). An evaluation of policy frameworks for addressing ethical considerations in learning analytics. *Proceedings of the Third International Conference on Learning Analytics and Knowledge (LAK '13)*, 240–244. <https://doi.org/10.1145/2460296.2460344>
- Tan, L. Y., Hu, S., Yeo, D. J., & Cheong, K. H. (2025). Artificial intelligence-enabled adaptive learning platforms: A review. *Computers and Education: Artificial Intelligence*, 9, Article 100429. <https://doi.org/10.1016/j.caeai.2025.100429>
- Zhu, H., Sun, Y., & Yang, J. (2025). Towards responsible artificial intelligence in education: A systematic review on identifying and mitigating ethical risks. *Humanities and Social Sciences Communications*, 12, Article 1111. <https://doi.org/10.1057/s41599-025-05252-6>