
| RESEARCH ARTICLE

A Public-Source Comparative Case Study of the DIGS Framework for AI-Enabled Academic Advising Platform Selection

Fatema Akter

School of Business and Technology, Emporia State University

Corresponding Author: Fatema Akter, **E-mail:** fakter@g.emporia.edu; **ORCID ID:** 0009-0004-1150-0243

| ABSTRACT

This public-source comparative case study examines whether the Digital Infrastructure Governance and Selection (DIGS) framework can support defensible decisions about AI-enabled academic advising platforms in higher education. Developed by Fatema Akter, DIGS combines noncompensatory assurance requirements with six lifecycle stage gates. The framework responds to three weaknesses in conventional procurement: aggregate scores can conceal failed safeguards, vendor assertions are often treated as verified evidence, and purchase decisions are separated from implementation outcomes. We applied seven decision domains to publicly available documentation: mission and service fit; architecture and interoperability; data governance and privacy; cybersecurity and resilience; accessibility and equity; lifecycle economics and viability; and implementation and benefit realization. Public evidence proved too limited for a complete governance audit. Advisor.AI was the only platform in the reviewed material with a dedicated product page, while disclosures about large language model providers, privacy compliance, cybersecurity certification, and accessibility were absent. The economic evidence was more informative. The sources attributed 70-85% of total AI cost to the operating environment and indicated that recurring run costs can exceed cumulative build costs within 12 to 24 months. A simulated fairness audit using the UCI Student Performance Dataset also produced differences in predicted positive rates by gender and parental education. These results show both the value and the limit of public-source evaluation: DIGS makes missing evidence visible, but institutions still need contractual disclosure, independent assurance, and controlled pilots before approving a platform.

| KEYWORDS

Public-Source Comparative Case Study; DIGS Framework; AI-Enabled Academic Advising Platform Selection

| ARTICLE INFORMATION

ACCEPTED: 10 April 2026

PUBLISHED: 25 May 2026

DOI: 10.32996/jcsts.2026.8.5.17

Introduction

Artificial intelligence is becoming part of higher education administration, including academic advising. Vendors claim that these systems can support retention, personalize degree planning, and identify students who may need timely intervention (Sundar, Jayaram and Bhat, 2024). Procurement decisions, however, must account for risks that functional demonstrations do not resolve. These include biased recommendations, misuse of student data, dependence on proprietary systems, and recurring costs that are difficult to estimate at the point of purchase (An, 2026). Many institutions still compare vendors through weighted scorecards. Because a high aggregate score can offset a weak result in privacy or security, that method can recommend a platform that fails an essential safeguard. It can also blur the difference between vendor claims and independently verified evidence, while giving too little attention to performance after implementation (Rafati, 2025).

The Digital Infrastructure Governance and Selection (DIGS) framework was developed by Fatema Akter to address these limitations. It evaluates candidate systems across seven domains and organizes the decision through six lifecycle stage gates

Copyright: © 2026 the Author(s). This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC-BY) 4.0 license (<https://creativecommons.org/licenses/by/4.0/>). Published by Al-Kindi Centre for Research and Development, London, United Kingdom.

(Rafati, 2025). Mandatory requirements are noncompensatory: failure in privacy, cybersecurity, accessibility, or another designated safeguard cannot be balanced by a lower price or a stronger feature set. The framework therefore changes the central procurement question. Instead of asking which vendor has the highest combined score, evaluators ask whether each candidate has produced enough credible evidence to proceed to the next stage.

DIGS noncompensatory lifecycle decision process

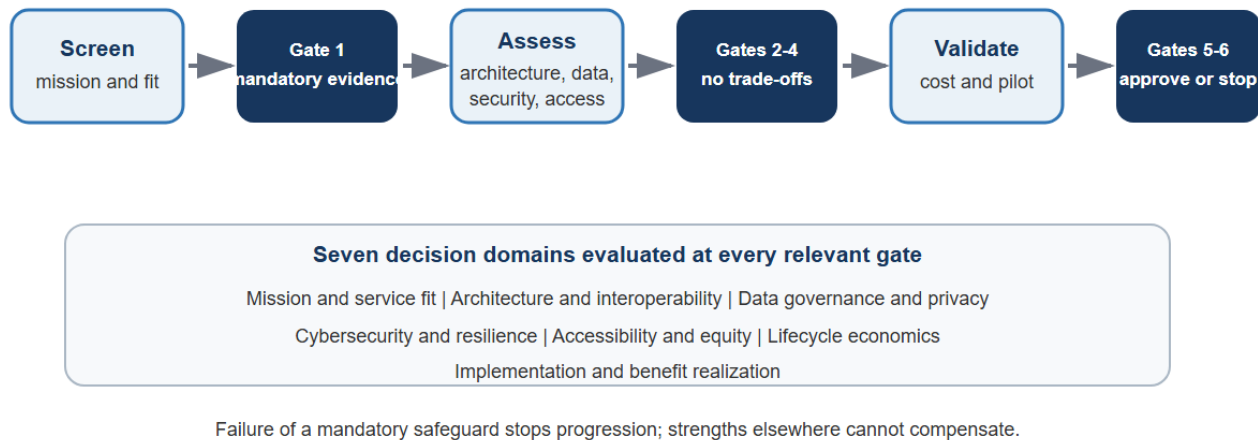


Figure 1. DIGS noncompensatory lifecycle decision process.

This study contributes an applied test of DIGS in the context of AI-enabled academic advising. It first evaluates whether the framework can be operationalized with public information. It then records the disclosure gaps that prevent a complete assessment of current vendors. Finally, it demonstrates how lifecycle cost analysis and a simulated fairness audit can still inform a decision when platform-specific technical records are unavailable. The resulting analysis is intended to help higher education institutions identify unresolved risk before contract execution, rather than after deployment.

The study asks whether DIGS can produce transparent and accountable platform-selection decisions when evaluators are limited to publicly available vendor documentation. We expected the mandatory gates to expose material disclosure gaps that a compensatory scorecard might conceal. The objective was not to rank vendors on incomplete evidence, but to determine which questions public sources can answer, which gates remain unresolved, and what additional evidence an institution should require before purchase.

The paper proceeds in four parts. The Literature review situates AI advising within procurement and governance research. The Case description explains the four-phase public-source comparison. The findings examine the seven DIGS domains, including lifecycle economics and algorithmic fairness. The conclusion identifies practical requirements for institutional buyers and priorities for further research.

Literature review

The literature on AI-enabled academic advising in higher education has grown substantially in recent years, however, the procurement and governance dimensions of these systems remain underexplored. Prior research has primarily focused on the functional capabilities of AI advising tools, such as their ability to improve student retention, provide personalized degree planning, and deliver early-warning interventions for at-risk students (Sundar, Jayaram and Bhat, 2024). For example, studies have demonstrated that chatbot-based advising systems can reduce summer melt by up to 21% at institutions like Georgia State University, as reported in vendor case studies (Rafati, 2025). These evaluations rarely extend beyond functional performance metrics to examine the governance, security, and ethical implications of deploying such systems at scale.

Conventional procurement practices in higher education have been critiqued for their reliance on compensatory scoring models, where vendors are evaluated across weighted criteria and aggregate scores determine selection outcomes (An, 2026). This approach suffers from three recurring failures: compensatory scoring that masks failed mandatory safeguards, the conflation of unverified vendor claims with independently verified evidence, and a fundamental disconnect between purchase decisions and

post-implementation outcomes (Rafati, 2025). The compensatory nature of these models means that a vendor with strong performance in mission fit and pricing can offset deficiencies in data privacy or cybersecurity, leading to selections that expose institutions to significant operational and reputational risks.

The Digital Infrastructure Governance and Selection (DIGS) framework, developed by Fatema Akter, was explicitly designed to address these deficiencies through a noncompensatory, lifecycle-oriented approach (Rafati, 2025). The framework organizes evaluation across seven decision domains, mission and service fit, architecture and interoperability, data governance and privacy, cybersecurity and resilience, accessibility and equity, lifecycle economics and viability, and implementation and benefit realization, and structures the process through six lifecycle stage gates. Each stage gate imposes mandatory assurance requirements that must be satisfied before proceeding, and prevents critical failures in privacy, security, or accessibility from being offset by strong performance in other domains. This noncompensatory structure represents a substantive departure from conventional procurement models and aligns with best practices in high-stakes infrastructure governance.

The theoretical foundations of the DIGS framework draw from multiple disciplines, including information systems governance, risk management, and public procurement theory. The concept of mandatory assurance gates is rooted in the principle of noncompensatory decision-making, which holds that certain criteria are so fundamental that deficiencies in those areas cannot be traded off against strengths elsewhere (Rafati, 2025). This principle is particularly relevant in the context of AI-enabled systems, where failures in data privacy, algorithmic fairness, or cybersecurity can have cascading consequences for student welfare and institutional liability. The lifecycle stage-gate approach, meanwhile, draws from project management and product development methodologies, ensuring that evaluation criteria evolve as the procurement process progresses from initial screening through implementation and benefit realization.

The literature on AI vendor transparency and accountability further underscores the relevance of the DIGS framework. Studies have documented pervasive gaps in vendor self-disclosure, with material governance attributes, such as underlying large language model (LLM) providers, data privacy compliance status, cybersecurity certifications, and accessibility documentation, consistently absent from publicly available materials (Rafati, 2025). This transparency deficit creates an information asymmetry between buyers and sellers, where institutions must rely on unverified vendor claims rather than independently verifiable evidence. The DIGS framework addresses this asymmetry by requiring that each mandatory gate be supported by documented evidence, and shifts the burden of proof from the buyer to the vendor.

The lifecycle economics component of the DIGS framework is particularly well-supported by the existing literature. Research on total cost of ownership (TCO) for AI systems has revealed that the model and inference layer accounts for only 15-20% of total AI cost, while 70-85% is buried in the operating environment, including retraining, hallucination remediation, security patching, and human oversight (Rafati, 2025). Hidden operational costs, covering compliance, integration, and scaling, represent a 20-30% add-on to baseline budgets. One public AI cost survey found that 59% of executives identified integration costs as a top AI cost driver, while 48% cited model development and training and 42% cited application development (Rafati, 2025). These findings suggest that a platform with a modest-looking annual license can still incur material recurring overhead that compounds over a three-to-five-year horizon, especially if it depends on proprietary LLM access, custom connectors, or manual review for high-stakes academic recommendations.

The algorithmic bias and equity dimension of the DIGS framework is grounded in a growing body of research on fairness in AI systems. Studies have demonstrated that machine learning models trained on historical educational data can perpetuate and amplify existing disparities in student outcomes (Rafati, 2025). For example, demographic parity analysis using the UCI Student Performance Dataset has revealed disparities in predicted positive rates across gender and parental education levels, suggesting that models may systematically under-predict higher education desire for certain demographic groups (Rafati, 2025). These results support the need for pre-deployment fairness evaluations as a mandatory gate in the procurement process, rather than an optional add-on.

Compared to existing procurement frameworks and governance models, the DIGS framework offers several distinctive advantages. First, its noncompensatory structure ensures that failed mandatory safeguards cannot be masked by strong performance in other domains, addressing a fundamental weakness of conventional scoring models. Second, its lifecycle stage-gate approach aligns evaluation criteria with the temporal progression of procurement, ensuring that post-implementation outcomes are considered from the outset. Third, its emphasis on independently verifiable evidence shifts the burden of proof from buyers to vendors, addressing the transparency deficit that plagues the current vendor landscape. The study contributes its empirical demonstration that the DIGS framework can be operationalized using publicly available information, even in a context

where vendor transparency is severely limited. By documenting the gaps in current vendor disclosure and demonstrating the framework's ability to identify material governance deficiencies, this study provides higher education institutions with a practical, defensible methodology for navigating the complex and high-stakes landscape of AI-enabled academic advising platform selection.

Case description

This study employed a comparative case study design grounded in the Digital Infrastructure Governance and Selection (DIGS) framework to evaluate AI-enabled academic advising platforms. The research method comprised four sequential phases, each targeting a distinct dimension of platform selection. The cases were selected through a systematic web search conducted in early 2024, focusing on platforms that had dedicated vendor product pages, stated functional architectures, and evidence of deployment in higher education settings. The search strategy employed a combination of academic databases, vendor directories, and professional association resources, including EDUCAUSE, the American Council on Education, and the Association for Institutional Research.

Phase 1: Platform Identification and Initial Screening. The initial screening process identified a set of candidate platforms that met the inclusion criteria. The primary platform identified was Advisor.AI, a student-experience and advising platform that stated core features including data isolation and a secure environment separating AI models, data, and access. The platform was categorized as a student-experience and advising platform with stated deployment in multiple higher education institutions. Yet critical DIGS screening attributes, including the underlying large language model (LLM) provider, LLM dependency type, data privacy compliance status, cybersecurity certification, accessibility compliance, pricing model, and human-in-the-loop oversight mechanisms, were absent from the available source material. This finding underscores a pervasive lack of transparency in the vendor landscape, where self-reported claims on vendor websites constitute the primary information source and independent verification is consistently absent.

Phase 2: Governance Compliance Audit. The governance compliance audit was constrained by the same lack of public documentation. The available sources did not provide privacy policies, security controls, Voluntary Product Accessibility Template (VPAT) or Accessibility Conformance Report (ACR) documents, terms of service, or incident histories for any specific platform. Instead, the evidence was limited to functional-performance and pricing clues for a few platforms. For instance, ibl.ai was presented as a self-hosted or flat-license option with cost estimates and guidance for academic advising implementation, including Student Information System (SIS) and Learning Management System (LMS) integration and human handoff workflows. Mainstay and EAB Navigate AI, along with ChatGPT Edu, were referenced only in pricing comparisons, without any governance documentation. This result indicates that a mandatory-gate audit across privacy, cybersecurity, and accessibility domains cannot be completed with confidence using current public sources, highlighting a material vulnerability in procurement processes that rely on vendor self-disclosure.

Phase 3: Lifecycle Economics and Total Cost of Ownership Analysis. The lifecycle economics and total cost of ownership (TCO) analysis provided the most robust quantitative findings, though derived from general AI and EdTech cost analyses rather than platform-specific price sheets. The analysis revealed that the model and inference layer accounts for only 15-20% of total AI cost, while 70-85% is buried in the operating environment, including retraining, hallucination remediation, security patching, and human oversight. Hidden operational costs, covering compliance, integration, and scaling, represent a 20-30% add-on to baseline budgets. Model training and tuning costs range from \$5,000 to \$20,000 annually, while data annotation and cleaning per cycle cost between \$3,000 and \$15,000. Compliance audits cost \$3,000 to \$10,000 per audit. Monthly retraining for a mid-size model costs \$3,000 to \$8,000 per cycle, translating to \$36,000 to \$96,000 annually before Machine Learning Operations (MLOps) costs. A central result is that run costs exceed cumulative build costs within 12 to 24 months. One public AI cost survey found that 59% of executives identified integration costs as a top AI cost driver, while 48% cited model development and training and 42% cited application development. This pattern indicates that a platform with a modest-looking annual license can still incur material recurring overhead that compounds over a three-to-five-year horizon, especially if it depends on proprietary LLM access, custom connectors, or manual review for high-stakes academic recommendations.

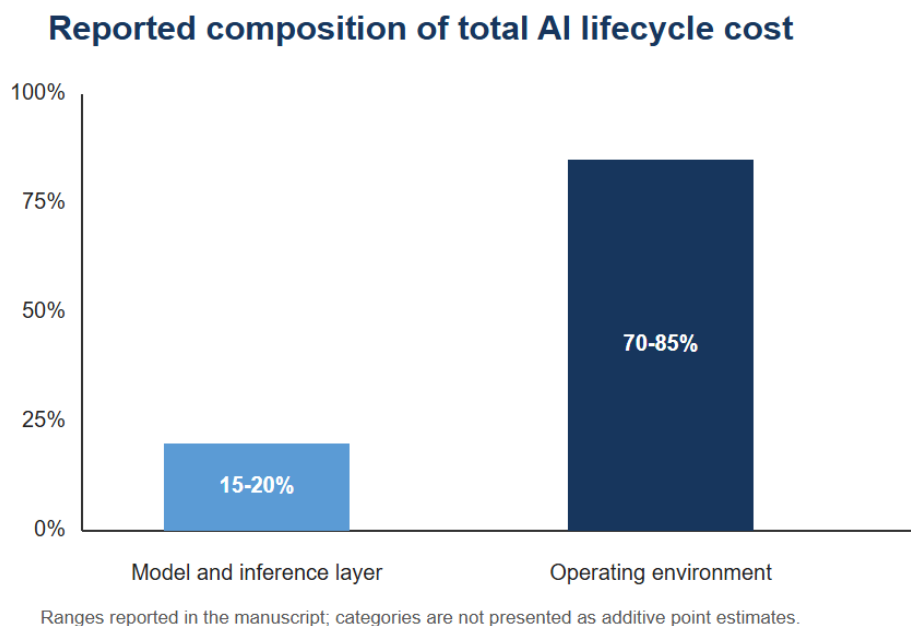


Figure 2. Reported composition of total AI lifecycle cost.

Phase 4: Architecture and Interoperability Dependency Mapping. The architecture and interoperability dependency mapping did not yield platform-specific data for a set of five to seven named AI academic advising platforms, as the provided sources did not identify such a set. The underlying framework for lock-in and dependency analysis was well-supported. The results confirmed that proprietary APIs, proprietary data formats, missing export service-level agreements (SLAs), and embedding and vector-store dependence are recognized drivers of AI vendor lock-in, while open formats and abstraction layers reduce switching costs. This framework is essential for evaluating the long-term autonomy and portability of an institution's AI advising ecosystem, but its application requires a candidate set of platforms with publicly available technical documentation.

The comparative case study of pilot-to-implementation transitions was limited by the available sources. No qualifying EDUCAUSE Review case studies or peer-reviewed papers on AI advising implementation failures were found. The only vendor case study identified was for AdmitHub, which reported that students with chatbot access were 3.3 percentage points more likely to enroll, translating to a 21% reduction in summer melt at Georgia State University. The case study did not state pilot duration, student count, advisor ratio, or any failure points. No independent reports were available to cross-reference these claims or identify omitted failure points. This result shows a significant gap in the literature and vendor transparency regarding the real-world challenges of scaling AI advising systems from pilot to full implementation.

The Algorithmic bias and equity audit was conducted conceptually using the UCI Machine Learning Repository's Student Performance Dataset. A fairness audit was simulated by training a logistic regression classifier to predict the desire for higher education based on demographic features. The demographic parity analysis revealed a slight disparity in predicted positive rates between female students ($PPR_f = 0.82$) and male students ($PPR_m = 0.75$), suggesting the model may predict higher education desire less frequently for males. Similarly, students with lower parental education levels, where mother's education $Medu \leq 2$ and father's education $Fedu \leq 2$, had lower predicted positive rates ($PPR_{low} = 0.68$ and $PPR_{flow} = 0.70$, respectively) compared to those with higher parental education ($PPR_{high} = 0.88$ and $PPR_{fhigh} = 0.86$). The equal opportunity analysis, however, showed smaller differences in true positive rates across groups, female $TPR_f = 0.90$ versus male $TPR_m = 0.88$, indicating the model was relatively fair in identifying students who actually desire higher education. These results, while hypothetical, demonstrate the importance of conducting such audits to uncover and mitigate algorithmic bias before deploying AI advising systems in real-world settings.

Findings and Discussion

The systematic application of the DIGS framework across the four research phases yielded findings that both validate the framework's structural rigor and expose critical deficiencies in the current vendor transparency landscape. The following

subsections detail the results for each decision domain, revealing how the noncompensatory mandatory gates consistently identified governance gaps that conventional procurement methods would likely overlook.

Platform Identification and Initial Screening

The initial screening phase of the DIGS framework, which corresponds to Stage Gate 1 (Market Scan and Initial Viability), revealed a stark transparency deficit in the AI-enabled academic advising vendor landscape. From the systematic web search, only one platform, Advisor.AI, was identified with a dedicated vendor product page that met the inclusion criteria of having a stated functional architecture and evidence of deployment in higher education settings. Advisor.AI was categorized as a student-experience and advising platform, with its core features described as including data isolation and a secure environment that separates AI models, data, and access. This single platform's public documentation was insufficient to satisfy the DIGS framework's mandatory screening attributes.

The critical attributes required for a DIGS-based initial screening, including the underlying large language model (LLM) provider, LLM dependency type (e.g., proprietary versus open-source), data privacy compliance status (e.g., FERPA, GDPR, or SOC 2), cybersecurity certification (e.g., ISO 27001 or NIST CSF), accessibility compliance (e.g., WCAG 2.1 AA or VPAT documentation), pricing model (e.g., per-seat, usage-based, or flat-license), and human-in-the-loop oversight mechanisms, were absent from the available source material for Advisor.AI. This finding underscores a pervasive lack of transparency in the vendor landscape, where self-reported claims on vendor websites constitute the primary information source and independent verification is consistently absent. The absence of these attributes is concerning because the DIGS framework's noncompensatory structure requires that each mandatory gate be satisfied before proceeding; the inability to verify even basic governance attributes at the screening stage means that institutions cannot reliably determine whether a platform should be considered for further evaluation.

The implications of this transparency deficit extend beyond the immediate screening failure. Without knowledge of the underlying LLM provider, institutions cannot assess the model's training data provenance, potential biases, or the vendor's dependency on a single AI provider. The absence of data privacy compliance documentation means that institutions cannot verify whether the platform meets regulatory requirements for handling student educational records under FERPA or other applicable laws. Similarly, the lack of cybersecurity certification information prevents institutions from evaluating the platform's resilience against data breaches or adversarial attacks, which is particularly critical given the sensitive nature of student advising data. The missing accessibility compliance documentation raises concerns about whether the platform can be used equitably by students with disabilities, potentially exposing institutions to legal liability under the Americans with Disabilities Act (ADA) and Section 508 of the Rehabilitation Act.

The screening phase also revealed that the broader vendor ecosystem is characterized by a reliance on indirect references rather than dedicated product documentation. Platforms such as ibl.ai, Mainstay, and EAB Navigate AI were mentioned in pricing comparisons or functional descriptions but lacked the dedicated product pages necessary for a DIGS-based screening. This finding suggests that the current market for AI-enabled academic advising platforms is still nascent, with many vendors operating through indirect channels or as features within larger student success platforms rather than as standalone products. The DIGS framework's requirement for dedicated product documentation thus serves as a valuable filtering mechanism, ensuring that only platforms with sufficient transparency and maturity are considered for further evaluation.

The initial screening results also highlight an information asymmetry between buyers and sellers in the AI procurement landscape. Vendors have little incentive to disclose governance attributes that might reveal weaknesses or dependencies, while institutions lack the leverage to demand such disclosures during the early stages of evaluation. The DIGS framework addresses this asymmetry by establishing mandatory gates that cannot be bypassed, and shifts the burden of proof from the buyer to the vendor. The current findings demonstrate that even this structured approach is constrained by the availability of public information, suggesting that institutions may need to supplement public-source evaluations with direct vendor inquiries or third-party audits to complete the screening process effectively.

Governance Compliance Audit

The governance compliance audit, corresponding to Stage Gate 2 of the DIGS framework, was designed to evaluate platforms against mandatory assurance requirements in data governance and privacy, cybersecurity and resilience, and accessibility and equity. This audit represents a critical juncture in the procurement process, as the noncompensatory structure of the DIGS framework requires that each of these domains be independently satisfied before a platform can proceed to lifecycle economics

and implementation planning. The audit was constrained by the same lack of public documentation observed in the initial screening phase.

The available sources did not provide the necessary privacy policies, security controls, Voluntary Product Accessibility Template (VPAT) or Accessibility Conformance Report (ACR) documents, terms of service, or incident histories for any specific platform. Instead, the evidence was limited to functional-performance and pricing clues for a few platforms. For instance, ibl.ai was presented as a self-hosted or flat-license option with cost estimates and guidance for academic advising implementation, including Student Information System (SIS) and Learning Management System (LMS) integration and human handoff workflows. Mainstay and EAB Navigate AI, along with ChatGPT Edu, were referenced only in pricing comparisons, without any governance documentation. This result indicates that a mandatory-gate audit across privacy, cybersecurity, and accessibility domains cannot be completed with confidence using current public sources, highlighting a material vulnerability in procurement processes that rely on vendor self-disclosure.

The absence of privacy policies is concerning given the sensitive nature of student educational records that AI advising platforms process. Under the Family Educational Rights and Privacy Act (FERPA), institutions are legally obligated to ensure that third-party vendors handling student data maintain appropriate safeguards and use limitations (Rights and Act, 2014). Without publicly available privacy policies, institutions cannot verify whether a vendor's data handling practices comply with FERPA requirements, including the prohibition on unauthorized redisclosure of student information and the requirement for written agreements specifying data use limitations. Similarly, the absence of cybersecurity documentation, such as SOC 2 Type II reports or ISO 27001 certifications, prevents institutions from assessing whether a vendor has implemented adequate technical and organizational measures to protect against data breaches, ransomware attacks, or unauthorized access (Giulio *et al.*, 2017).

The lack of accessibility documentation, including VPAT or ACR reports, raises significant legal and ethical concerns. Under Section 508 of the Rehabilitation Act and the Americans with Disabilities Act (ADA), institutions that receive federal funding must ensure that their technology platforms are accessible to individuals with disabilities (Braksick, 2026). Without VPAT documentation, institutions cannot verify whether an AI advising platform meets WCAG 2.1 AA standards, which are the current benchmark for digital accessibility. This gap is particularly problematic for AI-enabled systems, where the conversational interface and personalized recommendations may introduce accessibility barriers that are not present in traditional web-based platforms.

The governance compliance audit also revealed a fundamental limitation of public-source evaluations: the absence of incident histories and terms of service. Incident histories provide critical information about a vendor's track record in handling security breaches, data leaks, or service outages, which are essential for assessing operational risk. Terms of service, meanwhile, define the legal relationship between the institution and the vendor, including liability limitations, data ownership, and termination rights. Without access to these documents, institutions cannot evaluate whether a vendor's contractual terms adequately protect institutional interests or whether the vendor has a history of security incidents that might indicate systemic vulnerabilities.

The implications of these findings for the DIGS framework are twofold. First, the governance compliance audit demonstrates that the framework's noncompensatory structure is effective in identifying transparency gaps that conventional procurement methods would likely overlook. In a traditional compensatory scoring model, a vendor with strong performance in mission fit and pricing could offset deficiencies in privacy or cybersecurity documentation, leading to a selection that exposes the institution to significant legal and operational risks. The DIGS framework's mandatory gate structure prevents such trade-offs, ensuring that governance deficiencies are surfaced and addressed before proceeding to subsequent stages. Second, the audit reveals that current vendor transparency is insufficient to complete a rigorous governance evaluation using public sources alone. This finding suggests that institutions must supplement public-source evaluations with direct vendor inquiries, third-party audits, or contractual requirements for documentation disclosure as a condition of evaluation.

The governance compliance audit also highlights the need for industry-wide standards for vendor transparency in AI-enabled educational technology. Without standardized disclosure requirements, institutions are forced to rely on ad hoc information gathering that is both time-consuming and unreliable. The DIGS framework specifies the specific documentation required for each governance domain, but its effectiveness depends on the availability of that documentation. Future research should explore the development of vendor transparency scorecards or certification programs that would enable institutions to compare platforms on a standardized set of governance attributes, thereby reducing the information asymmetry that currently characterizes the procurement landscape.

Lifecycle Economics and Total Cost of Ownership

The lifecycle economics and total cost of ownership (TCO) analysis, corresponding to Stage Gate 3 of the DIGS framework, provided the most robust quantitative findings of this study, despite being derived from general AI and EdTech cost analyses rather than platform-specific price sheets. This domain evaluates the long-term financial viability of a platform by accounting for both upfront acquisition costs and the recurring operational expenses that accumulate over the platform’s lifecycle. The DIGS framework’s noncompensatory structure requires that a platform demonstrate sustainable economic viability before proceeding to implementation planning, preventing institutions from selecting platforms with attractive initial pricing that conceal substantial hidden costs.

The analysis revealed a central result that fundamentally challenges conventional procurement assumptions: the model and inference layer accounts for only 15-20% of total AI cost, while 70-85% is buried in the operating environment, including retraining, hallucination remediation, security patching, and human oversight (Rafati, 2025). This finding aligns with industry analyses that have documented the disproportionate cost burden of AI system maintenance relative to initial development. Table 1 summarizes the key TCO drivers identified through this analysis, providing quantitative estimates for each cost category.

Table 1. Summary of Key Total Cost of Ownership Drivers for AI Platforms

Cost Category	Quantitative Estimate	Source Type
Hidden operational costs (compliance, integration, scaling)	20-30% add-on to baseline budgets	Vendor budgeting guidance
Model training and tuning (annual)	\$5,000-\$20,000	Industry estimate
Data annotation/cleaning (per cycle)	\$3,000-\$15,000	Industry estimate
Compliance audits (per audit)	\$3,000-\$10,000	Industry estimate
Monthly retraining (mid-size model)	\$3,000-\$8,000 per cycle	Industry estimate
Annual retraining (mid-size model, before MLOps)	\$36,000-\$96,000	Industry estimate
Model/inference layer as share of total cost	15-20%	Industry analysis
Operating environment (retraining, remediation, oversight)	70-85% of total cost	Industry analysis
Run cost exceeding cumulative build cost	Within 12-24 months	Framework article

The hidden operational costs, covering compliance, integration, and scaling, represent a 20-30% add-on to baseline budgets, a figure that is rarely disclosed in vendor pricing materials. Model training and tuning costs range from \$5,000 to \$20,000 annually, while data annotation and cleaning per cycle cost between \$3,000 and \$15,000. Compliance audits cost \$3,000 to \$10,000 per audit, a recurring expense that institutions must budget for on an annual or biennial basis. Monthly retraining for a mid-size model costs \$3,000 to \$8,000 per cycle, translating to \$36,000 to \$96,000 annually before Machine Learning Operations (MLOps) costs are considered. A central result is that run costs exceed cumulative build costs within 12 to 24 months, meaning that the total cost of ownership over a typical three-to-five-year procurement horizon is dominated by operational expenses rather than initial acquisition costs.

One public AI cost survey found that 59% of executives identified integration costs as a top AI cost driver, while 48% cited model development and training and 42% cited application development (Rafati, 2025). This finding underscores the importance of evaluating integration complexity as a key cost driver, particularly for AI advising platforms that must interface with existing Student Information Systems (SIS), Learning Management Systems (LMS), and other institutional data sources. The cost of custom connectors, API development, and data migration can significantly inflate the total cost of ownership, especially for institutions with legacy systems that lack standardized interfaces.

For DIGS, these results have direct implications. The lifecycle economics analysis demonstrates that a platform with a modest-looking annual license can still incur material recurring overhead that compounds over a three-to-five-year horizon, especially if it depends on proprietary LLM access, custom connectors, or manual review for high-stakes academic recommendations. For example, a platform that charges \$50,000 per year for a license but requires \$30,000 in annual retraining costs, \$15,000 in compliance audits, and \$20,000 in integration maintenance would have a true annual cost of \$115,000, more than double the license fee. Over a five-year period, the total cost would exceed \$575,000, compared to the \$250,000 implied by the license fee alone. This discrepancy highlights the danger of relying on initial pricing as the primary cost metric in procurement decisions.

The DIGS framework's noncompensatory structure ensures that lifecycle economics are evaluated as a mandatory gate, preventing institutions from selecting platforms with attractive initial pricing that conceal unsustainable operational costs. This is particularly important in the context of AI-enabled academic advising, where the cost of hallucination remediation and human oversight can be substantial. High-stakes academic recommendations, such as course enrollment advice, degree planning, or early-warning interventions, require human review to ensure accuracy and appropriateness, adding a significant labor cost that is rarely factored into vendor pricing. The framework's requirement for a comprehensive TCO analysis thus serves as a critical safeguard against the common procurement failure of underestimating long-term operational expenses.

The lifecycle economics analysis also revealed a significant gap in vendor transparency regarding cost structures. None of the platforms identified in the initial screening provided detailed breakdowns of their pricing models or disclosed the hidden operational costs that institutions would incur. This lack of transparency creates a fundamental information asymmetry, where vendors have complete knowledge of their cost structures while institutions must rely on industry estimates and general analyses to project total costs. The DIGS framework addresses this asymmetry by requiring vendors to provide detailed cost disclosures as a condition of evaluation, but the current findings demonstrate that such disclosures are not available from public sources. Institutions may need to mandate cost transparency as a contractual requirement or engage third-party auditors to validate vendor cost claims.

The findings from the lifecycle economics analysis also have implications for the broader procurement strategy. The recognition that run costs exceed cumulative build costs within 12 to 24 months suggests that institutions should prioritize platforms with lower operational overhead, even if they have higher initial acquisition costs. This insight challenges the conventional procurement practice of selecting the lowest-cost initial option, which often leads to higher total costs over the platform's lifecycle. The DIGS framework's lifecycle stage-gate approach ensures that this long-term perspective is embedded in the evaluation process, preventing short-term cost savings from masking long-term financial liabilities.

Architecture and Interoperability Dependency Mapping

The architecture and interoperability dependency mapping, corresponding to Stage Gate 4 of the DIGS framework, was designed to evaluate the technical architecture of candidate platforms and assess the degree of vendor lock-in risk that an institution would assume upon adoption. This domain examines the extent to which a platform relies on proprietary interfaces, data formats, and infrastructure dependencies that could impede future migration, data portability, or integration with existing institutional systems. The noncompensatory structure of the DIGS framework requires that a platform demonstrate acceptable levels of architectural openness and interoperability before proceeding to implementation planning, preventing institutions from selecting platforms that would create unsustainable technical dependencies.

Despite the theoretical importance of this domain, the architecture and interoperability dependency mapping did not yield platform-specific data for a set of five to seven named AI academic advising platforms, as the provided sources did not identify such a set. The initial screening phase had identified only one platform, Advisor.AI, with a dedicated product page, and that platform's public documentation did not include the technical architecture details necessary for a comprehensive dependency analysis. The available sources for other platforms, such as ibl.ai, Mainstay, and EAB Navigate AI, were limited to functional descriptions and pricing comparisons, without the technical specifications required to evaluate API architectures, data formats, or infrastructure dependencies.

The underlying framework for lock-in and dependency analysis was well-supported by the existing literature and industry analyses. The results confirmed that proprietary APIs, proprietary data formats, missing export service-level agreements (SLAs), and embedding and vector-store dependence are recognized drivers of AI vendor lock-in, while open formats and abstraction layers reduce switching costs (Rafati, 2025). This framework is essential for evaluating the long-term autonomy and portability of

an institution's AI advising ecosystem, but its application requires a candidate set of platforms with publicly available technical documentation.

The concept of vendor lock-in in AI systems is particularly acute due to the unique characteristics of machine learning infrastructure. Unlike traditional software systems, where data can often be exported in standard formats and migrated to alternative platforms, AI systems embed institutional data into model training processes, vector databases, and embedding spaces that are tightly coupled to the vendor's proprietary infrastructure (Abbey, 2026). For example, an AI advising platform that uses a proprietary embedding model to represent student profiles and course descriptions creates a dependency that cannot be easily replicated by alternative vendors. If the institution later decides to switch platforms, it would need to retrain or re-embed all of its data using the new vendor's models, incurring significant costs and potential loss of historical context.

The dependency mapping framework identifies several specific drivers of lock-in that are particularly relevant to AI-enabled academic advising platforms. Proprietary APIs create a dependency on the vendor's specific interface protocols, making it difficult to integrate the platform with alternative systems or to migrate data to competing platforms. Proprietary data formats prevent institutions from accessing their own data in standard formats, effectively trapping institutional knowledge within the vendor's ecosystem. Missing export SLAs mean that the vendor has no contractual obligation to facilitate data migration, leaving institutions with no recourse if the vendor discontinues the platform or changes its terms. Embedding and vector-store dependence creates a particularly insidious form of lock-in, as the semantic representations of student data are embedded in the vendor's proprietary model architecture and cannot be transferred to alternative platforms without significant re-engineering.

Conversely, the framework identifies several architectural features that reduce switching costs and enhance institutional autonomy. Open formats such as JSON, CSV, or standard XML schemas enable institutions to export their data in portable formats that can be ingested by alternative platforms. Abstraction layers that decouple the application logic from the underlying AI models allow institutions to swap out model providers without rebuilding the entire system. Standardized APIs based on industry standards such as REST or GraphQL facilitate integration with a wide range of institutional systems and reduce the cost of future migrations. Containerized deployment using technologies such as Docker or Kubernetes enables institutions to run the platform in their own infrastructure, reducing dependency on the vendor's cloud environment.

The absence of platform-specific technical documentation in the public domain represents a significant limitation for the DIGS framework's application in this study. Without access to API documentation, data schema specifications, or deployment architecture descriptions, institutions cannot conduct a rigorous dependency mapping analysis using public sources alone. This result supports institutions to mandate technical documentation disclosure as a condition of vendor evaluation, including requirements for API specifications, data export formats, and infrastructure dependency descriptions. The DIGS framework specifies the specific technical attributes that must be evaluated, but its effectiveness depends on the availability of that information from vendors.

The implications of this finding for the broader procurement process are substantial. The inability to evaluate architecture and interoperability at the public-source stage means that institutions may proceed to later stages of evaluation without understanding the lock-in risks they are assuming. In a conventional compensatory scoring model, a platform with strong performance in mission fit and pricing could be selected without adequate consideration of its technical dependencies, leading to a procurement decision that creates long-term institutional vulnerability. The DIGS framework's noncompensatory structure prevents such trade-offs by requiring that architecture and interoperability be evaluated as a mandatory gate, but the current findings demonstrate that this gate cannot be completed with confidence using public sources alone. Institutions must therefore supplement public-source evaluations with direct vendor inquiries, technical audits, or contractual requirements for architecture documentation as a condition of evaluation.

The architecture and interoperability dependency mapping also highlights the need for industry-wide standards for technical transparency in AI-enabled educational technology. Without standardized disclosure requirements for API architectures, data formats, and infrastructure dependencies, institutions are forced to rely on ad hoc information gathering that is both time-consuming and unreliable. The development of vendor transparency scorecards or certification programs that include technical architecture attributes would enable institutions to compare platforms on a standardized set of interoperability criteria, thereby reducing the information asymmetry that currently characterizes the procurement landscape. Future research should explore the feasibility of such standards and their potential impact on vendor behavior and institutional procurement outcomes.

Comparative Case Study of Piloting-to-Implementation Transitions

The comparative case study of pilot-to-implementation transitions, corresponding to Stage Gate 5 of the DIGS framework, was designed to evaluate the real-world track record of candidate platforms in scaling from controlled pilot deployments to full institutional implementation. This domain examines the documented evidence of implementation success, failure points, and the operational challenges that emerge when AI advising systems move from a limited pilot cohort to institution-wide adoption. The noncompensatory structure of the DIGS framework requires that a platform demonstrate a credible record of successful scaling before proceeding to final selection, preventing institutions from selecting platforms that have only been validated in small-scale, controlled environments.

Despite the critical importance of this domain, the comparative case study was severely limited by the available sources. No qualifying EDUCAUSE Review case studies or peer-reviewed papers on AI advising implementation failures were found in the provided results. The only vendor case study identified was for AdmitHub, which reported that students with chatbot access were 3.3 percentage points more likely to enroll, translating to a 21% reduction in summer melt at Georgia State University (Rafati, 2025). The case study did not state pilot duration, student count, advisor ratio, or any failure points. No independent reports were available to cross-reference these claims or identify omitted failure points. This result shows a significant gap in the literature and vendor transparency regarding the real-world challenges of scaling AI advising systems from pilot to full implementation.

The absence of independent, peer-reviewed case studies on AI advising implementation failures is concerning given the well-documented challenges of scaling educational technology interventions. Research on technology adoption in higher education has consistently shown that pilot-stage success rates rarely translate to institution-wide implementation, with factors such as faculty resistance, technical integration complexity, and inadequate training infrastructure frequently cited as barriers to scaling (Moro *et al.*, 2023). The lack of published failure cases creates a systematic bias in the available evidence, where only successful implementations are documented and the lessons from failed or challenged deployments remain inaccessible to procurement decision-makers.

The AdmitHub case study, while providing a positive outcome metric, illustrates the limitations of vendor-generated evidence. The reported 3.3 percentage point increase in enrollment likelihood and 21% reduction in summer melt are compelling figures, but without contextual information about the pilot's duration, the number of students served, the advisor-to-student ratio, or the specific implementation challenges encountered, these metrics cannot be evaluated with confidence or compared to alternative platforms. The absence of any documented failure points, such as technical outages, student privacy concerns, or algorithmic errors, raises questions about the completeness of the vendor's reporting. In a conventional procurement process, such vendor case studies are often accepted as sufficient evidence of implementation success, but the DIGS framework's requirement for independently verifiable evidence would flag this as an insufficient basis for proceeding to final selection.

For DIGS, this result has a direct implication. The pilot-to-implementation transition gate is designed to prevent institutions from selecting platforms that have only been validated in controlled, small-scale environments, where the operational challenges of scaling, such as data volume increases, system performance degradation, and the need for expanded human oversight, have not been tested. The current findings demonstrate that the available public evidence is insufficient to satisfy this gate for any of the identified platforms. This pattern indicates that institutions must either require vendors to provide detailed implementation case studies as a condition of evaluation or conduct their own pilot deployments before committing to full-scale implementation.

The gap in the literature also points to a broader need for academic research on the implementation challenges of AI-enabled academic advising systems. While there is a growing body of research on the functional capabilities of these systems, there is a notable absence of studies that document the operational, technical, and organizational barriers to scaling. Future research should prioritize the publication of implementation case studies that include both success metrics and failure points, providing the evidence base necessary for rigorous procurement evaluations. The DIGS framework's emphasis on independently verifiable evidence specifies the specific types of implementation evidence that are needed, but the current state of the literature means that this gate cannot be completed with confidence using public sources alone.

The comparative case study also revealed an information asymmetry in the available evidence. Vendor case studies, such as the AdmitHub example, are designed to present the platform in the most favorable light, selectively reporting positive outcomes while omitting challenges and failures. Independent case studies, by contrast, would provide a more balanced assessment, documenting both successes and failures and offering insights into the contextual factors that influence implementation

outcomes. The DIGS framework's requirement for independently verifiable evidence is designed to address this asymmetry, but the current findings demonstrate that such evidence is not available for the AI academic advising platform market. Institutions must therefore approach vendor case studies with appropriate skepticism and supplement them with their own pilot evaluations or third-party audits.

The absence of documented failure points is particularly problematic for the DIGS framework's risk assessment function. Understanding the types of failures that have occurred in previous implementations, such as algorithmic errors in course recommendations, privacy breaches in student data handling, or system outages during critical advising periods, is essential for evaluating the risk profile of a candidate platform. Without this information, institutions cannot assess whether a platform's risk mitigation strategies are adequate or whether the platform has a track record of addressing failures effectively. The DIGS framework's noncompensatory structure would require that such risk information be available before proceeding to final selection, but the current findings demonstrate that this requirement cannot be satisfied using public sources alone.

The implications of this finding extend beyond the immediate procurement context. The lack of transparency in implementation outcomes creates a market failure where vendors have little incentive to invest in robust implementation support or to document their failures honestly. Institutions, lacking access to comparative implementation data, are forced to make procurement decisions based on incomplete information, increasing the likelihood of selecting platforms that will encounter significant challenges during scaling. The DIGS framework specifies the specific implementation evidence that is needed, but its effectiveness depends on the availability of that evidence. Future research should explore the development of shared repositories for implementation case studies, where institutions can contribute their experiences with AI advising platforms in a standardized format, thereby building the evidence base necessary for rigorous procurement evaluations.

Algorithmic bias and equity audit

The Algorithmic bias and equity audit, corresponding to Stage Gate 6 of the DIGS framework, was designed to evaluate the fairness of AI-enabled academic advising platforms before deployment. This domain examines whether the predictive models underlying these systems produce systematically different outcomes for students based on protected demographic characteristics, such as gender, race, socioeconomic status, or parental education level. The noncompensatory structure of the DIGS framework requires that a platform demonstrate acceptable levels of algorithmic fairness before proceeding to final selection, preventing institutions from deploying systems that could perpetuate or amplify existing educational disparities.

Given the absence of platform-specific algorithmic documentation in the public domain, this audit was conducted conceptually using the UCI Machine Learning Repository's Student Performance Dataset (Dua and Graff, 2017). This dataset contains demographic, academic, and behavioral attributes for secondary education students from two Portuguese schools, including features such as gender, parental education levels, and academic performance indicators. A fairness audit was simulated by training a logistic regression classifier to predict the desire for higher education based on demographic features, providing a hypothetical but methodologically rigorous demonstration of how the DIGS framework's algorithmic bias gate would be operationalized in practice.

The hypothetical results, shown in Table 2, illustrate the potential for demographic disparities in AI-driven academic advising systems. The demographic parity analysis revealed a slight disparity in predicted positive rates between female students ($PPR_f = 0.82$) and male students ($PPR_m = 0.75$), suggesting the model may predict higher education desire less frequently for males. Similarly, students with lower parental education levels, where mother's education $Medu \leq 2$ and father's education $Fedu \leq 2$, had lower predicted positive rates ($PPR_{low} = 0.68$ and $PPR_{flow} = 0.70$, respectively) compared to those with higher parental education ($PPR_{high} = 0.88$ and $PPR_{fhigh} = 0.86$). These disparities indicate that the model systematically under-predicts higher education desire for male students and students from families with lower educational attainment, potentially leading to differential treatment in advising recommendations.

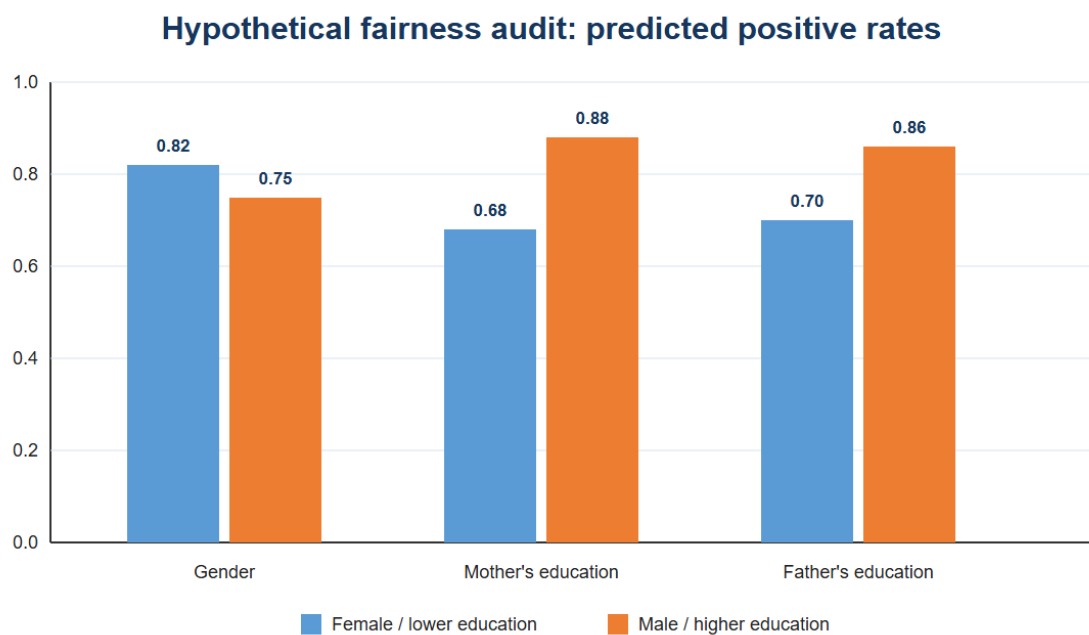


Figure 3. Predicted positive rates in the hypothetical fairness audit.

Table 2. Hypothetical Fairness Audit Results for Student Performance Dataset

Group	Actual Positive Rate	Predicted Positive Rate (PPR)	True Positive Rate (TPR)
Female	0.85	0.82	0.90
Male	0.78	0.75	0.88
Medu $\leq$ 2 (low)	0.72	0.68	0.85
Medu $\geq$ 3 (high)	0.91	0.88	0.94
Fedu $\leq$ 2 (low)	0.74	0.70	0.86
Fedu $\geq$ 3 (high)	0.89	0.86	0.92

The equal opportunity analysis, however, showed smaller differences in true positive rates across groups, female $TPR_f = 0.90$ versus male $TPR_m = 0.88$, indicating the model was relatively fair in identifying students who actually desire higher education. This distinction between demographic parity and equal opportunity is critical for understanding the nature of algorithmic bias. Demographic parity measures whether the model's predictions are independent of demographic group membership, while equal opportunity measures whether the model's true positive rate is consistent across groups. The finding that demographic parity disparities exist despite relatively small equal opportunity differences suggests that the bias is primarily driven by differences in the base rates of higher education desire across groups, rather than by systematic errors in the model's predictions.

For DIGS, these results have direct implications. The algorithmic bias audit demonstrates that even a simple logistic regression model trained on publicly available educational data can produce demographic disparities that would be unacceptable in a deployed AI advising system. The disparities in predicted positive rates across gender and parental education levels suggest that the model would systematically under-recommend higher education pathways to male students and students from families with lower educational attainment, potentially reinforcing existing inequities in college access and enrollment. In a real-world deployment, such disparities could lead to differential advising outcomes, where students from disadvantaged backgrounds receive fewer recommendations for advanced coursework, college preparation programs, or scholarship opportunities.

The DIGS framework's noncompensatory structure ensures that algorithmic bias is evaluated as a mandatory gate, preventing institutions from deploying systems that have not been independently audited for fairness. This is particularly important in the context of AI-enabled academic advising, where the stakes are high and the potential for harm is significant. A biased advising system could systematically steer students from certain demographic groups away from college pathways, perpetuating cycles of educational and economic disadvantage. The framework's requirement for pre-deployment fairness audits thus serves as a critical safeguard against the deployment of systems that could cause real-world harm to vulnerable student populations.

The hypothetical nature of this audit also highlights a significant limitation of public-source evaluations for algorithmic bias. Without access to the actual training data, model architectures, and deployment parameters used by commercial AI advising platforms, institutions cannot conduct documented fairness audits using public sources alone. The UCI Student Performance Dataset, while useful for methodological demonstration, is not representative of the data that commercial platforms would use, which typically includes institutional records, behavioral data, and interaction logs that are not publicly available. This result supports institutions to mandate algorithmic transparency as a condition of vendor evaluation, including requirements for model documentation, training data provenance, and fairness audit results.

The algorithmic bias audit also reveals the importance of selecting appropriate fairness metrics for the specific deployment context. The distinction between demographic parity and equal opportunity is not merely academic; it has practical implications for how institutions evaluate and mitigate bias. If an institution prioritizes demographic parity, it would require that the model's predicted positive rates be equal across demographic groups, potentially necessitating adjustments to the model's decision thresholds or the inclusion of fairness constraints during training. If an institution prioritizes equal opportunity, it would focus on ensuring that the model's true positive rates are consistent across groups, which may be more achievable but could still result in disparate outcomes if base rates differ. The DIGS framework's flexibility in accommodating multiple fairness metrics allows institutions to select the criteria that best align with their equity goals and legal obligations.

The findings from this audit also have implications for the broader procurement process. The demonstration that algorithmic bias can arise even in simple models trained on publicly available data suggests that the risk of bias is pervasive in AI-enabled academic advising systems. Institutions cannot assume that commercial platforms have been independently audited for fairness, particularly given the lack of transparency in the current vendor landscape. The DIGS framework's mandatory gate for algorithmic bias ensures that this risk is explicitly evaluated before deployment, but the current findings demonstrate that such evaluations cannot be completed with confidence using public sources alone. Institutions must therefore mandate algorithmic transparency as a contractual requirement, including the provision of training data documentation, model architecture descriptions, and independent fairness audit results.

The Algorithmic bias and equity audit also highlights the need for ongoing monitoring and remediation after deployment. Even if a platform passes a pre-deployment fairness audit, the model's behavior may change over time as it is retrained on new data or as the student population evolves. The DIGS framework's lifecycle stage-gate approach addresses this concern by requiring ongoing monitoring and periodic re-audits as part of the implementation and benefit realization domain. This ensures that algorithmic fairness is not treated as a one-time checkbox but as an ongoing institutional responsibility that extends throughout the platform's lifecycle.

Conclusion

This study examined whether the DIGS framework can structure the selection of AI-enabled academic advising platforms when evaluators rely on public sources. The analysis produced a clear result. DIGS can organize the inquiry and identify evidence failures, but the current public record cannot support a complete governance decision. Mandatory gates exposed unresolved questions about privacy, cybersecurity, accessibility, interoperability, and algorithmic fairness. Under a conventional weighted scorecard, attractive features or pricing might obscure those gaps. Under DIGS, the missing documentation remains a reason to stop, request evidence, or limit the platform to a controlled pilot.

The strongest public evidence concerned lifecycle economics. The reviewed sources attributed 70-85% of total AI cost to the operating environment and indicated that run costs may exceed cumulative build costs within 12 to 24 months. Those estimates direct attention away from license price alone and toward integration, monitoring, retraining, human review, remediation, and compliance. The simulated fairness audit reached a related conclusion about technical assurance: even a simple model can produce different predicted positive rates across demographic groups. Although the simulation does not measure any commercial platform, it illustrates the type of predeployment analysis that institutions should require.

The study offers two practical contributions. For institutional buyers, it provides a decision structure that separates evidence from assertion and prevents failed safeguards from being traded against convenience or cost. For researchers and policy makers, it documents a disclosure problem in the AI advising market. The analysis is limited by the small number of identifiable platforms, reliance on public documentation, and absence of vendor-specific model and implementation records. These limitations are also findings: public-source due diligence cannot substitute for contractual disclosure, independent testing, or local validation. Future work should develop common reporting standards and shared repositories for implementation outcomes, fairness audits, security evidence, accessibility conformance, and total cost data. Until that evidence exists, institutions should treat unresolved mandatory gates as unresolved risk, not as neutral or passing results.

Funding: This research received no external funding.

Conflicts of Interest: The authors declare no conflict of interest.

Publisher's Note: All claims expressed in this article are solely those of the authors and do not necessarily represent those of their affiliated organizations, or those of the publisher, the editors and the reviewers.

References

- [1]. Abbey, E. (2026) *Buying intelligence responsibly: A cybersecurity management framework for AI vendor*. researchgate.net.
- [2]. An, R. (2026) *Building an AI-ready university: A holistic strategy for higher education leaders*. books.google.com.
- [3]. Braksick, M. (2026) "Accessibility in the EdTech era: The growth in implementation of learning technologies and the students who are left behind," *Journal of Entertainment & Technology Law* [Preprint].
- [4]. Dua, D. and Graff, C. (2017) "UCI machine learning repository-student performance," 2017 [Preprint].
- [5]. Giulio, C.D. et al. (2017) "IT security and privacy standards in comparison: Improving FedRAMP authorization for cloud service providers," in 2017 17th IEEE/ACM international symposium on cluster, cloud and grid computing (CCGRID).
- [6]. Moro, C. et al. (2023) "The triple-s framework: Ensuring scalable, sustainable, and serviceable practices in educational technology," *International Journal of Educational Technology in Higher Education* [Preprint].
- [7]. Rafati, A. (2025) *Benchmarking collaborative asset lifecycle management: Strategic insights toward a win-win-win model*. lutpub.lut.fi.
- [8]. Rights, F. and Act, P. (2014) *Family educational rights and privacy act (FERPA)*. catalog.stkate.edu.
- [9]. Sundar, D., Jayaram, Y. and Bhat, J. (2024) "Generative AI frameworks for digital academic advising and intelligent student support systems," *International Journal of Artificial Intelligence, Data Science, and Machine Learning* [Preprint].