

EXPLORING AND RANKING AI-RELATED RISKS IN DOCTORAL STUDIES: AN EMPIRICAL AHP STUDY OF BUSINESS ADMINISTRATION RESEARCHERS IN VIETNAM

Vo Le Quynh Lam

Thu Dau Mot University

Email: lamvlq@tdmu.edu.vn

Received: 16 March 2026; Revised: 22 April 2026; Accepted: 10 May 2026

ABSTRACT

Amidst the swift transition to digital platforms, integrating artificial intelligence (AI) into academic inquiries offers substantial advantages but simultaneously introduces critical challenges for PhD candidates in the field of Business Administration. This research focuses on identifying and classifying the core risks linked to the adoption of AI, covering dimensions such as data integrity, ethical standards, technical vulnerabilities, and researcher proficiency. Utilizing the Analytic Hierarchy Process (AHP) framework, the investigation gathered insights from a sample of 175 PhD candidates across prominent Vietnamese Universities to evaluate the relative importance of these risk factors. The findings identify four major risk groups: (1) Data risks (weight: 0.56), encompassing incomplete data, biased data and data security concerns; (2) Ethical risks (0.26), involving challenges such as transparency, privacy violations, and AI decision-making bias; (3) Technical risks (0.12), including software errors and inaccurate AI models; and (4) Competence risks (0.06), which pertain to a lack of data analysis skills and difficulties in selecting appropriate AI tools. The research emphasizes the necessity of recognizing and controlling these potential threats, alongside offering strategic recommendations to reduce their impact and bolster AI-related expertise for PhD scholars.

Keywords: Analytic Hierarchy Process (AHP), artificial intelligence, Business Administration, doctoral students, risks in research.

1. INTRODUCTION

The confluence of rapid digital innovation and expansive data availability has rendered the fusion of technology with critical thinking indispensable in higher education and research [1]. In this vein, AI has evolved beyond its role as a simple utility to become a transformative element that radically reshapes the research frameworks used by doctoral researchers in Business Management. As García-Chitiva and Correa [1], contemporary education transcends rote knowledge dissemination, prioritizing instead the cultivation of analytical acumen and adaptive problem-solving skills that AI can augment but also potentially undermine if not managed judiciously.

This digital impetus has catalyzed advancements in academic inquiry, empowering PhD candidates with sophisticated technologies to expedite processes from data analysis to content generation [2]. Generative AI tools, such as ChatGPT, exemplify this potential by facilitating diverse outputs from textual syntheses to visual representations thereby streamlining tasks like literature synthesis and hypothesis formulation [3, 4]. However, this integration is not without peril: over-dependence on AI may erode independent thought, propagate inaccuracies, or

engender ethical dilemmas, as evidenced by empirical critiques highlighting risks to research authenticity [5].

Critically, extant literature reveals a conspicuous lacuna in addressing AI's holistic risks within the doctoral research lifecycle, often confining discourse to superficial benefits like personalized learning or administrative efficiency [4]. This oversight is particularly acute in emerging economies like Vietnam, where national AI strategies propel adoption yet grapple with governance gaps [6]. Gaps persist in quantifying AI's influence on tangible outcomes - such as thesis timelines, publication rates, and qualitative rigor - while qualitative explorations of user experiences in risk navigation remain underexplored. Addressing these voids is imperative: by delineating risks through empirical lenses, this study posits a framework that not only optimizes AI utility but also fortifies doctoral resilience, thereby elevating scholarly competitiveness in Business Administration amid global digital shifts.

2. LITERATURE REVIEW

2.1. Theoretical background

AI is fundamentally a key pillar of computer science, focused on developing systems capable of replicating human cognitive functions to support learning, environmental adjustment, and independent decision processes [7]. Propelled by surges in computational prowess, voluminous datasets, and advanced machine learning paradigms, AI's evolution - particularly through Generative AI (GenAI) models like large language models - has broadened its remit into intricate domains such as natural language processing and creative synthesis [3, 8]. In Vietnam, this trajectory aligns with policy-driven digital transformation, fostering AI assimilation across socio-economic spheres while mandating ethical oversight to avert systemic vulnerabilities [6, 9]. Such developments fundamentally reconfigure research ecosystems in Business Administration, where empirical, data-centric insights are paramount for strategic decision-making.

Theoretically, AI's affordances in research manifest through automation of laborious tasks, enabling efficient navigation of vast datasets and revelation of latent patterns that elude conventional methodologies [10, 11]. For instance, machine learning algorithms underpin predictive analytics in supply chain dynamics or market forecasting, bridging theoretical models with practical business exigencies. In academic praxis, AI expedites literature curation by interrogating expansive repositories, distilling insights, and even prototyping drafts, thereby amplifying efficiency [4, 12]. Collaborative synergies are further enhanced via AI-mediated platforms that democratize data sharing and interdisciplinary discourse [8, 13]. In Vietnam's nascent AI milieu, these capabilities resonate with national R&D imperatives, ostensibly elevating productivity and innovation in Business Administration doctoral pursuits [6]. However, such positive outlooks require cautious consideration: as noted by Zhu [14] in a comprehensive analysis of AI's influence on professional growth in academia, although AI promotes training based on specific competencies, it may worsen the gap between skills and industry needs if educational programs fail to keep pace with digital innovations.

2.2. Potential Threats of Integrating AI into Academic Study and Investigation

Notwithstanding AI's theoretical promise, its deployment in doctoral research engenders risks that, if unmitigated, imperil academic integrity and scholarly rigor. Data risks, for example, stem from inherent biases or incompleteness in inputs, precipitating erroneous inferences - a concern amplified in Vietnam by localized data privacy deficits and algorithmic prejudices that distort business-oriented analyses [3, 5]. Ethically, AI's opacity in decision pathways fosters plagiarism, attribution lapses, and integrity breaches, as Lierena-Izquierdo [15] critiques in a systematic review of AI ethics in academia, identifying transparency and bias as pivotal

challenges across disciplines. Technical vulnerabilities, including systemic failures and cyber threats, compound these issues, potentially compromising sensitive datasets in PhD inquiries [16]. Competence risks, meanwhile, erode foundational skills: over-reliance may stifle critical thinking and creativity, as warned by prior studies [4, 17], with empirical evidence linking AI dependency to diminished cognitive engagement in higher education.

These risks are not abstract; they intersect theory and practice profoundly. Alawamleh et al. [13], employing Interpretive Structural Modeling to dissect AI limitations in business, delineate 15 interlinked factors—such as trust deficits, ethical quandaries, and privacy erosions - that mirror the hierarchies uncovered in this study's AHP framework. Similarly, a decision matrix for prioritizing GenAI risks in higher education leverages AHP to rank 24 perils, underscoring data security and ethical biases as dominant, akin to Vietnamese doctoral contexts where resource constraints exacerbate such vulnerabilities [16]. Literature gaps abound: frameworks for risk classification in doctoral settings are scant, data quality's empirical impacts underexplored, ethical/security dilemmas under-theorized, critical thinking's erosion unquantified, and tailored management guidelines for Business Administration absent [2, 18, 19, 20]. This study contends that bridging these chasms through AHP-driven prioritization not only theorizes risks but operationalizes mitigation, fostering a balanced AI-research nexus.

I anchor this study within the Technology Acceptance Model (TAM) and Institutional Theory. Specifically, I frame AI risks not merely as technical glitches but as 'inhibitors' that increase the perceived risk of adoption, thereby disrupting the 'human-capital-augmented' research process by creating barriers to trust and academic integrity.

3. METHODOLOGY

To rigorously prioritize AI risks in doctoral Business Administration research, the current research utilizes the Analytic Hierarchy Process (AHP) framework. This multi-criteria approach integrates subjective evaluations with numerical weightings through a process of hierarchical structuring and pairwise analysis [16]. AHP's efficacy in risk assessment is evidenced in analogous contexts, such as prioritizing GenAI perils in higher education, where it elucidates inter-criteria dependencies.

Regarding the expert panel, participants were selected using a purposive sampling technique, focusing on senior doctoral candidates and faculty members with at least three years of experience in AI-integrated research to ensure the validity of the pairwise comparisons. Furthermore, while this study utilizes a standard AHP model, it explicitly acknowledges the assumption of criterion independence. Potential overlaps between categories, such as the intersection of data security and technical errors, were addressed during the pre-survey briefing to ensure participants evaluated each risk as a distinct entity, thereby minimizing multi-collinearity effects.

The methodology unfolds in two phases. Phase 1 entails a systematic literature review (2020-2024) across WoS, Scopus, and Google Scholar, yielding four principal risk clusters data, ethical, technical, and competence encompassing 12 sub-criteria. Phase 2 integrates qualitative expert consultations with quantitative surveys from 175 PhD students in Business Administration at Ho Chi Minh City's premier institutions: Industrial University, University of Economics, University of Economics and Law, and University of Finance and Marketing. While AHP traditionally employs a small expert panel, I justify this 'hybrid' approach by viewing 175 doctoral candidates as 'practitioner-experts'. They are the primary agents navigating AI integration in the research lifecycle. This larger sample provides a broader representative weight of "user-perceived" risks, which is crucial for informing university policy, though I acknowledge this may prioritize operational risks over systemic technical ones.

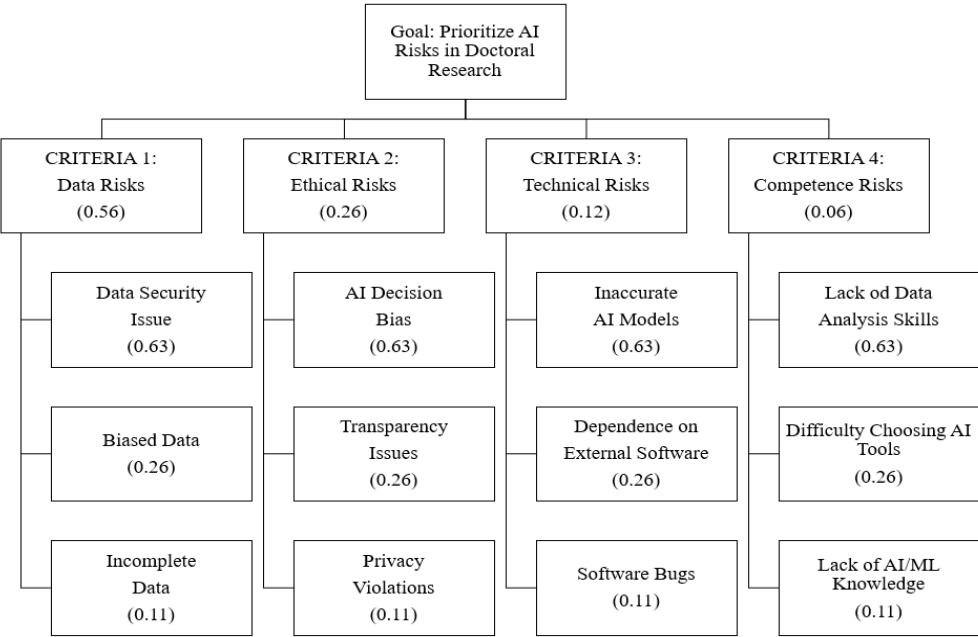


Figure 1. AHP Hierarchical Tree (Source: Compiled by the author)

Demographically, the sample comprises 65% males and 35% females; age distributions span 28-30 (30%), 31-40 (50%), and 41-55 (20%); occupations include lecturers (40%) and managers (60%); incomes range from 10-20 million VND (40%) to over 20-30 million VND (60%). This composition ensures representativeness, reflecting Vietnam's doctoral cohort's diversity.

Risk criteria are delineated as: Data risks (incomplete data, biased data, data security); Ethical risks (AI decision bias, transparency issues, privacy violations); Technical risks (software bugs, inaccurate AI models, dependence on external software); Competence risks (lack of data analysis skills, lack of AI knowledge, difficulty in tool selection). Pairwise evaluations employ Saaty’s 1-9 scale, with weight derived via the principal eigenvector and consistency verified ($CR < 0.1$), aligning with methodological rigor in empirical risk studies [16].

4. RESULTS AND DISCUSSION

The AHP analysis yields a prioritized hierarchy of AI risks, as encapsulated in Table 1.

Table 1. Comparison Matrix and AHP weight for Key Criteria

Priority	Risk Type	Preferred vector weight
1	Data risk	0.56
2	Ethical risks	0.26
3	Technical risks	0.12
4	Competence risks	0.06
n	4	
RI	0.9	
λ_{max}	4.18	
CI	0.06	
CR	0.07	< 0.1

(Source: Compiled by the author)

I have critically re-examined the low weighting assigned to competence risks (0.06). This finding, while seemingly contradictory to literature emphasizing AI-induced cognitive erosion, reflects a specific phenomenon among Vietnamese doctoral candidates: 'automation overconfidence.' It suggests that researchers may prioritize immediate operational efficiency over long-term analytical rigor. However, this undervaluation of competence risk is itself a significant risk, as a lack of deep critical engagement with AI outputs can lead to a decline in the qualitative rigor of doctoral dissertations, potentially compromising the depth of original contributions. I interpret the low weight of competence risks (0.06) through the lens of "automation bias" Doctoral students may overestimate their proficiency with AI, failing to recognize the subtle 'cognitive erosion' or loss of critical thinking often cited in recent literature. This disconnect suggests that researchers are more concerned with immediate data security than with the long-term degradation of their analytical skills. With $CR = 0.07 < 0.1$, the matrix exhibits robust consistency, affirming the reliability of doctoral perceptions. Data risks predominate at 0.56, a finding that resonates with global critiques: as Alawamleh et al. (2024) posit, data biases and insecurities undermine AI's foundational integrity, particularly in data-intensive Business Administration research where flawed inputs cascade into strategic misjudgments. The high priority of data risks (0.56) is contextualized by Vietnam's institutional voids in data governance. This finding aligns with the KPMG [9] report, which highlights that cyber vulnerabilities in Vietnamese academic institutions often lead to significant data breaches. Furthermore, the fact that AI bias outranks privacy suggests that local researchers prioritize functional accuracy over individual data confidentiality at this stage of AI implementation.

Sub-criteria breakdowns further illuminate these priorities (Tables 2 and 3). In technical risks, inaccurate AI models (0.63) eclipse software bugs (0.11) and external dependencies (0.26), underscoring the peril of unreliable outputs - a concern mirrored in Cao's [10] review of AI in business research, where algorithmic inaccuracies distort empirical validity. Within the data risk cluster, security concerns (0.63) represent the highest priority. This finding aligns with the Vietnamese context, where cyber vulnerabilities often lead to academic data breaches, as highlighted by KPMG [9]. The prioritization of AI bias (0.63) over privacy (0.11) reflects a cultural and stage-specific nuance in Vietnam. At this stage of AI implementation, researchers are more concerned with 'functional fairness' and the accuracy of results (to pass peer review) than with individual data confidentiality.

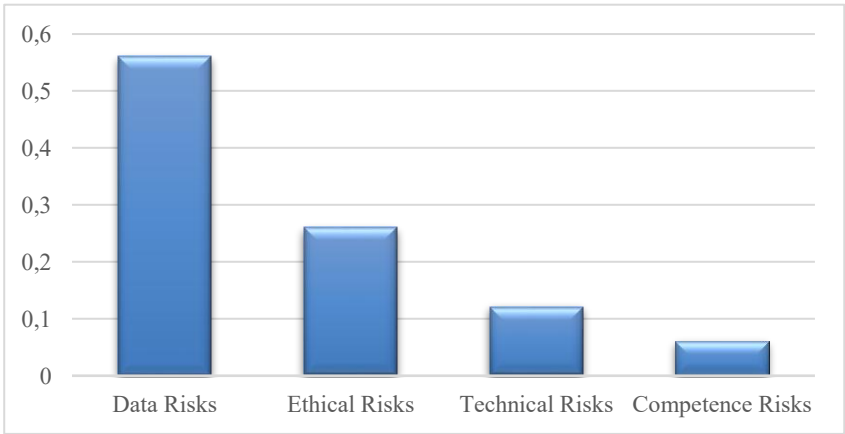


Figure 2. Primary AI Risk Comparison (Source: Compiled by the author)

To enhance comparative clarity, Figure 1 illustrates the hierarchical decomposition of AI risks, while Figure 2 visualizes the significant disparity between primary risk vectors, emphasizing the predominance of data-related concerns.

Table 2. AHP Detailed Breakdown for Technical and Data Risks

Technical risks	Software bugs	Inaccurate AI models	Dependence on external software	Preferred vector
Software bugs	1.00	0.20	0.33	0.11
Inaccurate AI models	5.00	1.00	3.00	0.63
Dependence on external software	3.00	0.33	1.00	0.26
Total	9.00	1.53	4.33	1.00
Data risks	Incomplete data	Biased data	Data security issues	Preferred vector
Incomplete data	1.00	0.33	0.20	0.11
Biased data	3.00	1.00	0.33	0.26
Data security issues	5.00	3.00	1.00	0.63
Total	9.00	4.33	1.53	1.00

(Source: Calculated by the author)

Ethical risks prioritize AI decision bias (0.63) over transparency (0.26) and privacy violations (0.11), a hierarchy that critiques AI's opaque "black box" nature, as Llerena-Izquierdo [15] argues, fostering unfair outcomes in research ethics. Competence risks emphasize data analysis skill deficits (0.63), surpassing tool selection difficulties (0.26) and AI knowledge gaps (0.11), echoing Wu's [21] call for curriculum reforms to bridge theory-practice divides in AI education.

Table 3. AHP Detailed Analysis for Ethical and Competence Risks

Ethical risks	AI decision bias	Transparency issues	Privacy violations	Preferred vector
AI decision bias	1.00	3.00	5.00	0.63
Transparency issues	0.33	1.00	3.00	0.26
Privacy violations	0.20	0.33	1.00	0.11
Total	1.53	4.33	9.00	1.00
Competence risks	Lack of data analysis skills	Lack of AI/machine learning knowledge	Difficulty in choosing AI tools	Preferred vector
Lack of data analysis skills	1.00	5.00	3.00	0.63
Lack of AI/machine learning knowledge	0.20	1.00	0.33	0.11
Difficulty in choosing AI tools	0.33	3.00	1.00	0.26
Total	1.53	9.00	4.33	1.00

(Source: Calculated by the author)

First, Data risks (0.56) emerged as the primary concern, which is contextualized by Vietnam's current institutional voids in data governance. The dominance of data security (0.63) within this cluster reflects localized cyber vulnerabilities. This finding aligns with the KPMG [9] report, which highlights that academic and research institutions in Vietnam are increasingly targeted by data breaches. As noted in the literature, while global concerns often focus on AI's

technical bugs, Vietnamese doctoral candidates prioritize the integrity and safety of their empirical data, seeing it as the most critical inhibitor to successful technology adoption.

Second, regarding Ethical risks (0.26), a provocative finding is that AI decision bias (0.63) significantly outranks privacy violations (0.11). While Western contexts often prioritize individual privacy, this cohort of Business Administration researchers focuses more on algorithmic transparency and fairness. I theorize that this reflects the current stage of AI implementation in Vietnam, where researchers are more concerned with "functional accuracy" and ensuring their methodologies are not distorted by biased AI outputs such as biased market forecasting—than with personal data confidentiality. This cultural nuance suggests that in Vietnamese academia, the validity of research outcomes is perceived as more ethically sensitive than data privacy.

Third, the low weight assigned to Competence risks (0.06) reveals a critical logical tension. Although literature warns of the "erosion of independent thought" due to AI, participants in this study seem to undervalue this threat. I argue that this disconnect is a clear manifestation of "automation bias" or "over-confidence bias," where doctoral students may overestimate their proficiency in AI tools while underestimating the subtle degradation of their own critical thinking skills. This finding underscores the need for "Human-in-the-loop" (HITL) requirements to preserve the qualitative rigor of doctoral dissertations.

5. CONCLUSION

The synthesized results from the AHP analysis indicate that data-related threats (0.56) are the primary concern, followed by ethical issues (0.26), technical vulnerabilities (0.12), and skill-based risks (0.06), compelling a reevaluation of AI's role in doctoral Business Administration. This hierarchy, bolstered by sub-criteria emphases (e.g., 0.63 on data security and AI bias), bridges theoretical vulnerabilities with practical imperatives, advocating robust governance frameworks to harness AI's potential while averting perils. To move beyond generic suggestions, I propose that institutions implement "AI-Verification Audit Trails" Since data risks dominate (0.56), doctoral candidates must maintain a log of AI-assisted data cleaning to ensure empirical validity. Furthermore, to address ethical risks (0.26), universities should align AI usage with strict 'Thesis Timelines,' ensuring that the speed of AI does not bypass the necessary period of human reflection, thereby safeguarding the intellectual integrity of the PhD process. To strengthen the coherence of this study, the identified quantitative weights serve as the direct foundation for the proposed framework. Specifically, the 0.56 weight for data risks necessitates the 'Data Verification Protocol' mentioned earlier, effectively bridging the gap between the vulnerabilities identified in the introduction and the practical imperatives of doctoral research in Vietnam. I propose concrete, granular interventions: first, mandatory AI ethics modules for PhD programs; second, 'Data Verification Protocols' to ensure the integrity of AI-generated results; and third, 'Human-in-the-loop' requirements to ensure that AI serves only as an auxiliary tool, preserving the researcher's critical judgment in complex tasks like market forecasting. Moving beyond generic recommendations, the proposed mitigation strategies directly address thesis timelines and research integrity. For instance, the 'Data Verification Protocol' is designed to prevent the cascading errors that occur when AI-generated data is unverified, thus safeguarding the empirical validity of the final dissertation. In conclusion, this study fills the research gap identified in the introduction by transforming qualitative AI concerns into a prioritized quantitative hierarchy. By linking the dominant Data risks (0.56) directly to mandatory verification protocols, the paper ensures that the proposed mitigation strategies are not merely generic suggestions but are empirically grounded interventions designed to safeguard the Vietnamese doctoral lifecycle. I propose specific interventions: (1) Mandatory AI Ethics Modules in doctoral curricula; (2) Data

Verification Protocols where candidates must document AI's role in data cleaning; and (3) a 'Human-in-the-loop' (HITL) requirement to ensure human oversight in critical business modeling like market forecasting.

Limitations include the study's sectoral (Business Administration) and geographic (Ho Chi Minh City) focus, subjective expert inputs, and confinement to four risk domains amid AI's dynamism. Future trajectories encompass cross-disciplinary expansions, methodological hybrids (e.g., AHP-DEMATEL), and probes into nascent risks like intellectual property, informed by evolving global discourses.

Acknowledgement: The author gratefully acknowledges the invaluable support and participation of the 175 PhD students from major universities in Ho Chi Minh City and the expert contributors, whose input greatly enriched this empirical study.

REFERENCES

- [1] M. d. P. Garcia-Chitiva and J. C. Correa, "Soft skills centrality in graduate studies offerings," *Studies in Higher Education*, vol. 49, no. 6, pp. 956-980, Jun. 2024, <https://doi.org/10.1080/03075079.2023.2254799>
- [2] E. De Nito, T. A. Rita Gentile, T. Köhler, M. Misuraca, and R. Reina, "E-learning experiences in tertiary education: patterns and trends in research over the last 20 years," *Studies in Higher Education*, vol. 48, no. 4, pp. 595-615, Apr. 2023, <https://doi.org/10.1080/03075079.2022.2153246>
- [3] A. E. Essien, O. Bukoye, C. O'Dea, and M. D. Kremantzis, "The influence of AI text generators on critical thinking skills in UK business schools," *Studies in Higher Education*, vol. 49, no. 5, pp. 865-882, May 2024, <https://doi.org/10.1080/03075079.2024.2316881>
- [4] M. Hosseini, M. Murad, and D. B. Resnik, "Benefits and risks of using AI agents in research," *Hastings Center Report*, vol. 56, no. 1, pp. 13-17, Jan.-Feb. 2026, <https://doi.org/10.1002/hast.70025>
- [5] M. Hooda, C. Rana, P. Goyal, S. Kadyan, and K. Devi, "Artificial intelligence for assessment and feedback to enhance student success in higher education," *Mathematical Problems in Engineering*, vol. 2022, Art. no. 5215722, pp. 1-19, Jun. 2022, <https://doi.org/10.1155/2022/5215722>
- [6] T. H. Bui, "The impact of artificial intelligence and digital economy on Vietnam's legal system," *Journal of Human Rights and Social Work*, vol. 7, no. 3, pp. 245-256, Sep. 2022, <https://doi.org/10.1007/s41134-022-00224-w>
- [7] T. H. Davenport and D. D. D'Ignazio, "How companies can prepare for the coming 'AI-first' world," *Strategy & Leadership*, vol. 51, no. 1, pp. 26-30, Jan. 2023, <https://doi.org/10.1108/SL-11-2022-0111>
- [8] C. Huff, "The promise and perils of using AI for research and writing," *Monitor on Psychology / APA Topics*, Oct. 2024. [Online]. Available: <https://www.apa.org/topics/artificial-intelligence-machine-learning/ai-research-writing>
- [9] KPMG in Vietnam, "Innovation and caution: A deep dive into generative AI risks for Vietnam's businesses," *Lexology*, Insights Report, Jan. 2024. [Online]. Available: <https://www.lexology.com/library/detail.aspx?g=cc8ce8da-b6c5-4445-af3a-ad75934440ed>

- [10] Z. Cao, M. Li, and P. A. Pavlou, "AI in business research," *Decision Sciences*, vol. 55, no. 5-6, pp. 518-532, Dec. 2024, doi: <https://doi.org/10.1111/deci.12655>
- [11] N. H. Tien, T. Le Nhut, and P. L. N. Thu, "Impact of AI on R&D performance of enterprises," *Cuestiones de Fisioterapia*, vol. 54, no. 5, pp. 438-450, Mar. 2025, doi: <https://doi.org/10.48047/pqcz9c71>
- [12] M. Guidoum and A. Elkhansa, "Challenges and risks of AI usage on students' academic performance: A case study at the University of Nizwa, Oman," *Creative Education*, vol. 15, no. 4, pp. 560-578, Apr. 2024, doi: <https://doi.org/10.4236/jilsa.2026.181003>
- [13] M. Alawamleh, N. Shammas, K. Alawamleh, and L. Bani Ismail, "Examining the limitations of artificial intelligence in business using Interpretive Structural Modelling," *Journal of Open Innovation: Technology, Market, and Complexity*, vol. 10, no. 3, Art. no. 100338, pp. 1-12, Sep. 2024, doi: <https://doi.org/10.1016/j.joitmc.2024.100338>.
- [14] H. Zhu, "The impact of generative AI use on graduate students' research competence: The mediating role of critical thinking and the moderating role of research self-efficacy," *Behavioral Sciences*, vol. 14, no. 3, Art. no. 204, pp. 1-15, Mar. 2024, doi: <https://doi.org/10.3390/bs16020304>.
- [15] J. Llerena-Izquierdo, "Ethics of the use of artificial intelligence in academia and research: The most relevant approaches, challenges and topics," *Informatics*, vol. 11, no. 4, Art. no. 71, pp. 1-18, Dec. 2024, doi: <https://doi.org/10.3390/informatics12040111>
- [16] M. Guillen-Aguinaga *et al.*, "Data quality in the age of AI: A review of governance, ethics, and the FAIR principles", *Data* vol. 10 no. 12, art. no. 201, Dec. 2025, doi: <https://doi.org/10.3390/data10120201>
- [17] T. Sari *et al.*, "Decision matrix for prioritizing generative AI risks in higher education," *Research Square*, Preprint, Jan. 2024, doi: <https://doi.org/10.21203/rs.3.rs-7494968/v1>
- [18] A. Calma and M. Davies, "Critical thinking in business education: Current outlook and future prospects," *Studies in Higher Education*, vol. 46, no. 11, pp. 2279-2295, Nov. 2021, doi: <https://doi.org/10.1080/03075079.2020.1716324>
- [19] A. Q. Nguyen, "Innovation in Vietnamese higher education for Society 5.0: Opportunities, challenges, and strategies," *International Journal of Research in Education and Science*, vol. 10, no. 4, pp. 750-760, 2024, doi: <https://doi.org/10.46328/ijres.3491>
- [20] T. Mazloun & B. Shahin, "AI and Research: Benefits and Challenges", *Clinical Oral Science and Dentistry*, vol. 6 no. 3, 2024, pp. 1-10. https://researchnovelty.com/management_research/article_pdf/1714981806Finalized%20Article_RA190424..pdf
- [21] Q. Wu, L. Chen, M. Chen, and Y. Huang, "Exploring the impact of artificial intelligence on business talent development in higher education: A systematic literature review and research agenda," *The International Journal of Management Education*, vol. 24, no. 1, Art. no. 101287, Mar. 2026, pp. 1-13, doi: <https://doi.org/10.1016/j.ijme.2025.101287>