

ARTIFICIAL INTELLIGENCE AND MANAGERIAL DECISION-MAKING: A SYSTEMATIC LITERATURE REVIEW OF ITS IMPLICATIONS FOR STRATEGIC MANAGEMENT IN ORGANIZATIONS

INTELIGENCIA ARTIFICIAL Y TOMA DE DECISIONES GERENCIALES: UNA
REVISIÓN SISTEMÁTICA DE LITERATURA SOBRE SUS IMPLICACIONES PARA LA
GESTIÓN ESTRATÉGICA EN LAS ORGANIZACIONES

Iván Miguel García López¹ y Jessica Nájera Ochoa²

SUMARIO: 1. Introduction, 2. Theoretical Framework, 2.1. Managerial Decision-Making Enhanced by Artificial Intelligence, 2.2. Organizational Governance of Artificial Intelligence, 2.3. Managerial Risk and Responsibility in Human-AI Decision Systems, 3. Method, 4. Results, 5. Discussion, 6. Conclusions, Sources of information

ABSTRACT

Artificial intelligence (AI) has been incorporated into business administration to support managerial decision-making; however, recent literature agrees that its strategic value lies not in full automation, but in augmenting human judgment, which introduces governance and responsibility challenges. The objective of this study is to systematically analyze how recent literature addresses the use of AI in managerial decision-making, integrating the constructs of decision augmentation, organizational AI governance, and managerial risk and responsibility. Methodologically, it is a systematic literature review applied to articles indexed in Scopus and Web of Science. The results show a predominance of decision-augmentation and procedural

RESUMEN

La inteligencia artificial (IA) se ha incorporado a la administración de empresas como soporte a la toma de decisiones gerenciales; no obstante, la literatura reciente coincide en que su valor estratégico no reside en la automatización plena, sino en el aumento del juicio humano, lo que introduce desafíos de gobernanza y responsabilidad. El objetivo de este estudio es analizar sistemáticamente cómo la literatura reciente aborda el uso de la IA en la toma de decisiones gerenciales, integrando los constructos de toma de decisiones gerenciales aumentada, gobernanza organizacional de la IA y riesgo y responsabilidad gerencial. Metodológicamente, se trata de una Revisión Sistemática de Literatura aplicada a artículos indexados en Scopus y

¹ Investigador en inteligencia artificial, modelos abiertos de lenguaje y gestión estratégica en instituciones de educación superior. Sus líneas de trabajo integran ciencias computacionales, derecho, ética y administración aplicada a entornos educativos.

² ULSA México. Doctora en Administración, Licenciada en Informática por el IPN y Doctora en Sostenibilidad por la UEMA. Cuenta con 28 años de experiencia en TIC, Administración y Sostenibilidad; Gerencia, Jefatura y Docencia en Posgrado.

governance approaches, as well as concerns about responsibility and human-AI decision reliance. As a contribution, the study offers an integrated managerial perspective to guide responsible organizational practices regarding AI.

KEYWORDS: artificial intelligence, managerial decision-making, organizational AI governance, managerial responsibility.

Web of Science. Los resultados evidencian un predominio de enfoques de aumento decisional y gobernanza procedimental, así como preocupaciones por la rendición de cuentas y la dependencia humano-IA. Como aporte, el estudio ofrece una lectura gerencial integrada para orientar prácticas organizacionales responsables de IA.

PALABRAS CLAVE: inteligencia artificial, toma de decisiones gerenciales, gobernanza organizacional de IA, responsabilidad gerencial.

1. Introduction

The incorporation of artificial intelligence into organizations has become one of the most relevant phenomena in business administration over the past decade, particularly due to its impact on the way managerial decisions are generated, evaluated, and executed (Ben-Michael et al., 2024). As algorithmic systems are integrated into strategic and operational processes, the literature has documented not only significant opportunities for improvement in efficiency and decision-making quality and new tensions related to organizational responsibility, control, and governance (Afroogh et al., 2024). This shift has moved the debate beyond a purely technological discussion toward a broader organizational and managerial perspective..

In this context, recent research agrees that the value of artificial intelligence for companies does not lie in the full

automation of decisions, but in its ability to enhance human judgment through advanced analytics, simulations, and informational support (Ameen et al., 2021). This approach to decision augmentation has been identified as the dominant paradigm in the literature, recognizing the bounded rationality of decision-makers and the need to preserve human agency in complex environments (Buijsman et al., 2025). However, the adoption of this paradigm raises fundamental questions regarding how IA use is governed, monitored and legitimized within organizations.

In parallel, the literature has begun to highlight that the effective implementation of AI in business administration depends to a large extent on the existence of organizational governance frameworks capable of translating ethical and strategic principles into concrete operational practices (OECD, 2025). Recent studies indicate that, without clear governance

mechanisms, AI can generate significant risks related to algorithmic opacity, loss of traceability, and dilution of managerial responsibility (Díaz-Rodríguez et al., 2023). These concerns have placed governance and risk management at the center of academic and practical debate.

Despite these advances, the field presents considerable fragmentation in terms of approaches, levels of analysis, and types of empirical evidence, which hinders an integrated understanding of AI's role in managerial decision-making (Duan et al., 2023). Many studies approach the phenomenon from isolated perspectives (technological, ethical, or strategic) without systematically articulating the relationship between decisional purpose, organizational governance and managerial risk (Silva, 2022). This dispersion limits the literature's ability to provide clear guidance to managers and organizational policymakers.

Against this backdrop, the aim of this study is to systematically analyze how recent literature conceptualizes and addresses the use of artificial intelligence in managerial decision-making, integrating three key constructs: augmented managerial decision-making, organizational AI governance, and managerial risk and responsibility. Unlike previous reviews focused on technical or ethical aspects in isolation, this paper proposes an integrated managerial reading (Chang et al., 2023) based on a systematic mapping of literature that classifies the existing evidence based on five research questions explicitly oriented to business administration. This approach allows not only to synthesize the state of

the art, but also to identify patterns, gaps, and tensions that guide future research and organizational practices. These elements are developed in the following sections.

2. Theoretical framework

2.1 Augmented Managerial Decision-Making through Artificial Intelligence

The integration of artificial intelligence into managerial decision-making has been predominantly conceptualized as a process of cognitive augmentation of the human decision-maker, rather than as a complete replacement of his or her judgment (Abhishek et al., 2025). From this perspective, AI acts as an amplifier of analytical capabilities, allowing managers to process massive volumes of information, identify complex patterns, and reduce the uncertainty inherent in contemporary organizational environments (AI-Kfairy et al., 2024). This approach aligns with the vision of decision-making as a socio-technical activity in which technology complements, but does not replace, the bounded rationality of the manager.

Recent literature highlights that AI systems generate managerial value when they are integrated into existing decision-making processes, respecting the organization's responsibility structures and strategic frameworks (Batool et al., 2024). Empirical studies show that the use of AI for decision augmentation improves the quality of decisions, increases the consistency of judgments, and facilitates the exploration of alternative strategic scenarios, particularly in contexts of high complexity and dynamism (Jiang et al., 2024). These findings reinforce the idea that AI's impact

is greatest when it is oriented towards informational and analytical support rather than full automation.

However, this construct also implies recognizing that decision augmentation introduces new challenges, such as the adequate calibration of trust and human-AI decision reliance, aspects that emerge across the study's RQs. In this sense, the augmentation approach establishes the conceptual framework from which both governance mechanisms and associated managerial risks are analyzed (Bacon, 2024), thereby connecting directly to the subsequent constructs.

2.2 Organizational Governance of AI in Managerial Contexts

The organizational governance of artificial intelligence has established itself as a central element to ensure that its use in managerial decision-making is consistent, responsible, and aligned with the organization's strategic objectives (Grimmelikhuijsen & Meijer, 2022). This construct is based on the premise that AI cannot be managed solely as a technological tool but requires governance frameworks that integrate clearly defined policies, processes, and roles within the organizational structure (Cheng & Wu, 2024).

Recent evidence shows a shift from approaches based solely on abstract ethical principles to procedural governance models, in which AI is regulated by operational protocols, internal standards, and control mechanisms embedded in existing management systems (Del Pozo et al., 2023). This type of governance makes

it possible to translate organizational values into concrete practices, facilitating the supervision of the use of AI and the mitigation of risks associated with its adoption (Ahmed Ali Linkon et al., 2024). Likewise, the literature underlines that the effectiveness of these frameworks depends on their ability to adapt to changing contexts and varying levels of digital maturity.

From a managerial perspective, AI governance not only serves a normative function but also an enabling one, by building organizational trust and legitimizing the use of algorithmic systems in critical decisions (Casey et al., 2019). This construct connects directly with RQs related to risk management and responsibility and establishes the conceptual bridge necessary to analyze how human-AI interaction is managed in practice, the central theme of the following construct (Ananny & Crawford, 2018).

2.3 Managerial Risk and Responsibility in Human-AI Decision Systems

The use of artificial intelligence in managerial decision-making introduces a specific set of organizational risks that transcend the technical aspects and are at the core of managerial responsibility (Bommasani et al., 2021). This construct is based on the idea that the adoption of AI redefines the boundaries of responsibility, authority, and control within organizations, generating tensions between efficiency, transparency, and accountability (Black et al., 2022).

The literature emphasizes that one of the main managerial risks associated with AI is algorithmic opacity, which hinders the

auditability of decisions and complicates the attribution of responsibilities in case of errors or adverse results (Ding et al., 2022). Recent studies warn that, without clear oversight and documentation mechanisms, AI systems can erode managers' ability to justify their decisions to internal and external stakeholders (Abunaser et al., 2025). This risk is intensified when AI is used in strategic or high-impact organizational decisions-making.

In addition, the risk of overconfidence and over-reliance on algorithmic systems emerges as a recurring concern, especially in contexts where AI is perceived as objective or infallible (Baathman, 2021). Research suggests that effectively managing these risks requires explicit human-in-the-loop approaches and trust calibration strategies, which preserve human agency and reinforce managerial responsibility (Li et al., 2025). Together, this construct provides the analytical framework to interpret the empirical results of the study and to inform the discussion on the practical implications of AI in business administration.

3. Method

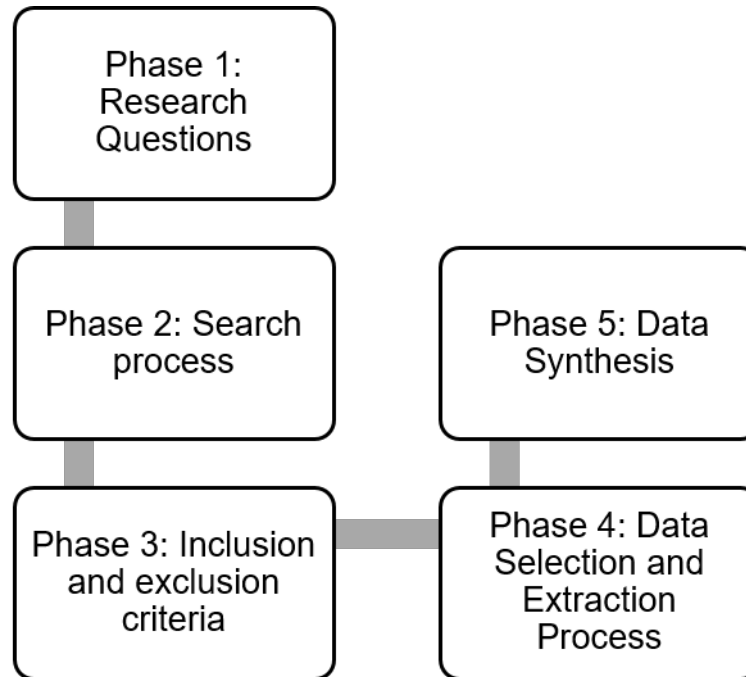
This study adopts a Systematic Literature Review (SLR) to identify, organize, and synthesize the scholarly evidence on artificial intelligence in managerial decision-making. The methodological design was structured to ensure transparency, replicability, and analytical rigor, following the PRISMA 2020 reporting framework and recent methodological guidance for systematic literature reviews in management research (Page et al., 2021). This approach is particularly appropriate

given the interdisciplinary nature of the topic, which spans strategic management, organizational governance, and socio-technical decision systems.

The primary purpose of the SLR is to identify, analyze, and integrate how recent literature conceptualizes and empirically addresses the use of artificial intelligence in managerial decision-making contexts, with a specific focus on three interrelated constructs: decision augmentation, organizational AI governance, and managerial risk and responsibility (Sauer & Seuring, 2023). Rather than concentrating on technical performance or isolated ethical considerations, the review adopts a managerial and organizational lens, enabling a systematic mapping of how AI is embedded in decision processes, governed within organizations, and associated with new forms of risk and responsibility. This framing ensures alignment between the methodological design and the theoretical objectives of the study.

The review process comprised five sequential phases (Paul et al., 2021): (1) formulation of research questions, (2) search strategy design and database selection, (3) definition of inclusion and exclusion criteria, (4) screening and study selection, and (5) data extraction and synthesis. These phases were designed to progressively refine the corpus while preserving conceptual breadth and methodological consistency, in accordance with recent guidance for transparent and replicable literature reviews. Fig. 1 summarizes the phases followed in the review process.

Fig. 1 Phases of the Systematic Literature Review



As an initial step, research questions (RQs) were formulated to provide analytical direction and ensure coherence between the review protocol and the study's objectives. These questions were explicitly designed to capture how the literature addresses (i) the primary managerial purposes of AI use, (ii) the types of organizational AI governance mechanisms employed, (iii) the management of human-AI decision reliance, (iv) the evidence of organizational value generated, and (v) the managerial risk domains highlighted. The formulation of these RQs reflects a deliberate effort to bridge fragmented streams of research and to structure the subsequent analysis in a way that supports an integrated managerial interpretation of the findings, which is presented in the following subsection corresponding to Phase 1 of the SLR.

Phase 1: Research Questions

Derived from the systematic review protocol adopted in this study, the first phase focused on the formulation of research questions (RQs) capable of capturing the most analytically relevant dimensions of artificial intelligence use in managerial decision-making contexts. Rather than prioritizing descriptive or purely bibliometric indicators, the RQs were designed to enable an in-depth examination of how artificial intelligence is conceptualized, governed, and problematized within organizational decision processes, with particular attention to managerial purposes, governance mechanisms, and risk domains. This orientation aligns with methodological recommendations that emphasize the role of research questions in structuring interpretative synthesis within Systematic

Literature Reviews in management research.

The development of the RQs was informed by an initial scoping of the literature, which revealed a high degree of fragmentation across technological, ethical, and strategic perspectives on AI, as well as a lack of integrative managerial frameworks. Building on these observations, the questions were formulated to explicitly address gaps related to the functional role of AI in managerial decision-making, the organizational arrangements governing its use, and the risks associated with human-AI interaction in decision systems. This approach ensured that the RQs were grounded in existing theoretical debates while remaining oriented toward unresolved issues in the field.

In line with the objectives of the study, the RQs were structured to correspond directly with the three core constructs guiding the analysis: decision augmentation, organizational AI governance, and managerial risk and responsibility. The possible answer categories associated with each question were derived deductively from the theoretical foundations of the study and inductively refined during the review process, allowing for both conceptual coherence and empirical sensitivity. This dual logic supports a systematic classification of the literature while preserving the flexibility required to capture emerging patterns.

Ultimately, the formulation of the research questions (Table 1) reflects the study's overarching motivation to advance a managerial understanding of artificial

intelligence beyond isolated technical or ethical considerations by examining how AI is actually embedded in organizational decision-making practices. By anchoring the SLR around these RQs, the review provides a structured foundation for the subsequent phases of article selection, coding, and analysis and enables a coherent synthesis of findings that directly inform both theory development and managerial practice.

These research questions defined the analytical dimensions used to examine the selected articles and structured the subsequent stages of the review process. In particular, the RQs guided the identification of conceptual and empirical intersections across studies, allowing the review to capture how artificial intelligence is simultaneously framed in terms of managerial purpose, organizational governance, and associated risks. This approach facilitated the detection of recurring patterns, tensions, and gaps in the literature, especially at the intersections between augmented decision-making practices, governance arrangements, and challenges related to responsibility, transparency, and managerial accountability. By explicitly focusing on these cross-cutting elements, the research questions ensured coherence between the theoretical framework and the empirical synthesis developed throughout the study.

Phase 2: Search Process

The document search process was conducted using Scopus and Web of Science (WoS), selected due to their extensive coverage of high-quality,

Table 1. Research Questions

Question	Possible responses
RQ1. What is the primary managerial purpose of AI use described in the study?	• Decision augmentation • Decision automation • Strategic analysis and planning • Operational optimization • Governance, compliance, and risk control (Acosta-Enriquez et al., 2025)
RQ2. What type of organizational AI governance is reported or implied?	• Structural governance (roles, oversight bodies, responsibility design) • Procedural governance (lifecycle controls, monitoring, evaluation routines) • Relational governance (stakeholder coordination, culture, human alignment) • Audit-based governance (internal algorithmic auditing end-to-end) • Principles-only governance (high-level ethics with limited operationalization) (Larsson, 2020)
RQ3. How is human-AI decision reliance managed (trust/overreliance/underreliance)?	• Trust calibration via explanations and confidence cues • Cognitive forcing / friction to reduce overreliance • Human-in-the-loop control (human final authority) • Algorithmic-in-the-loop / automated decisions with limited human review • No explicit reliance management reported (Zhai et al., 2024)
RQ4. What evidence is used to claim organizational value from AI (managerial outcomes)?	• Decision quality improvement (accuracy/consistency/robustness) • Speed and scalability of analysis (time-to-decision, throughput) • Risk reduction and control effectiveness (compliance, monitoring) • Workforce/HR productivity impacts (algorithmic management effects) • Conceptual or speculative value only (no measurable outcomes) (Mišić et al., 2025)
RQ5. What risk domain is foregrounded (managerial risk framing)?	• Responsibility and auditability risk • Governance implementation gaps (from principles to practice) • Trust/overreliance risk in decision-making • Organizational and stakeholder risk (context, effectiveness, legitimacy) • Workplace control and socio-technical risk (algorithmic management) (European Union, 2024)

peer-reviewed journals in the fields of management, information systems, and organizational studies. The use of these databases is consistent with best practices for systematic literature reviews, as they provide robust indexing standards and facilitate the retrieval of relevant and methodologically rigorous research. Employing both databases also reduced the risk of publication bias and enhanced the comprehensiveness of the search.

The search strategy was designed to identify studies addressing the use of artificial intelligence in managerial and organizational decision-making contexts. To this end, combinations of keywords related to artificial intelligence,

included to ensure the academic rigor and comparability of the evidence analyzed.

All searches were conducted on January 15, 2026, ensuring a clear temporal cutoff and enhancing the replicability of the review process. The search results obtained from Scopus and WoS were exported and consolidated for subsequent screening and selection procedures. Table 2 presents the search strings applied in Scopus and Web of Science in order to enhance methodological transparency and facilitate the replicability of the review process. Including these strings makes explicit how the corpus was operationally delimited and supports the rigor of the subsequent inclusion and exclusion stages.

Table 2. Search Strings

Search string in Scopus	Search string in WoS
TITLE-ABS-KEY("artificial intelligence" AND ("decision making" OR "decision support")) AND TITLE-ABS-KEY("managerial" OR "management" OR "strategic") AND TITLE-ABS-KEY(organization* OR business* OR firm* OR enterprise*) AND PUBYEAR > 2020 AND LIMIT-TO(DOCTYPE,"ar")	TS=("artificial intelligence" AND ("decision making" OR "decision support")) AND TS=("managerial" OR "management" OR "strategic") AND TS=(organization* OR business* OR firm* OR enterprise*) AND PY=(2021-2026) AND DT=(Article)

managerial decision-making, organizational governance, and risk or responsibility were applied to titles, abstracts, and author keywords. The temporal scope of the review was restricted to the period 2021–2025, reflecting the rapid acceleration of AI adoption in organizational settings and the corresponding growth of scholarly attention to its managerial implications. Only peer-reviewed journal articles were

Phase 3: Inclusion and Exclusion Criteria
Given the objective of analyzing how the recent literature addresses the use of artificial intelligence in managerial decision-making contexts, it was essential to define inclusion and exclusion criteria that ensured conceptual alignment, methodological rigor, and relevance to organizational and managerial research. Artificial intelligence encompasses a wide range of technologies and applications; therefore, the review

explicitly focused on studies that examine AI as part of organizational decision processes rather than as a purely technical artifact or isolated computational tool. This refinement was necessary to avoid conceptual dilution and to ensure that the selected corpus directly contributed to the study's analytical framework.

The inclusion criteria required that articles explicitly address the use of artificial intelligence in managerial, strategic, or organizational decision-making contexts, with clear relevance to at least one of the three core constructs guiding the review: decision augmentation, organizational AI governance, or managerial risk and responsibility. To ensure topical relevance, selected studies were required to include terms related to artificial intelligence and decision-making, governance, responsibility, risk, or management within their titles, abstracts, or author keywords. The temporal scope was limited to publications between January 2021 and December 2025, reflecting the accelerated adoption of AI in organizations and the corresponding emergence of managerial and governance-related debates in the literature.

Only peer-reviewed journal articles were included in the review to ensure academic quality, methodological consistency, and comparability of findings. This decision aligns with established SLR practices in management and information systems research, which emphasize the use of rigorously vetted sources when synthesizing evidence.

The exclusion criteria were applied to maintain analytical focus and methodological rigor. Book chapters, books, conference proceedings, editorials, and review articles were excluded, as these sources typically provide conceptual overviews or secondary syntheses rather than original empirical or theoretical contributions. Articles published outside the defined time window were also excluded to avoid outdated perspectives or speculative discussions not grounded in the current state of organizational AI adoption. Additionally, studies that addressed artificial intelligence exclusively from a technical, engineering, or algorithmic optimization perspective without explicit consideration of managerial decision-making or organizational implications were removed from the corpus. Together, these criteria ensured that the final dataset was both conceptually coherent and directly aligned with the research objectives.

Phase 4: Data Selection and Extraction Process

The initial search yielded 1256 articles from Scopus and 15 articles from Web of Science, resulting in a total of 1271 records. After removing 35 duplicate entries, the remaining articles were subjected to a systematic screening process based on titles, abstracts, and keywords. This step aimed to verify that artificial intelligence was addressed as a central element of managerial or organizational decision-making, rather than as a peripheral or purely technical topic.

During the screening phase, 427 articles were excluded because they did not focus on the use of AI in managerial decision-making

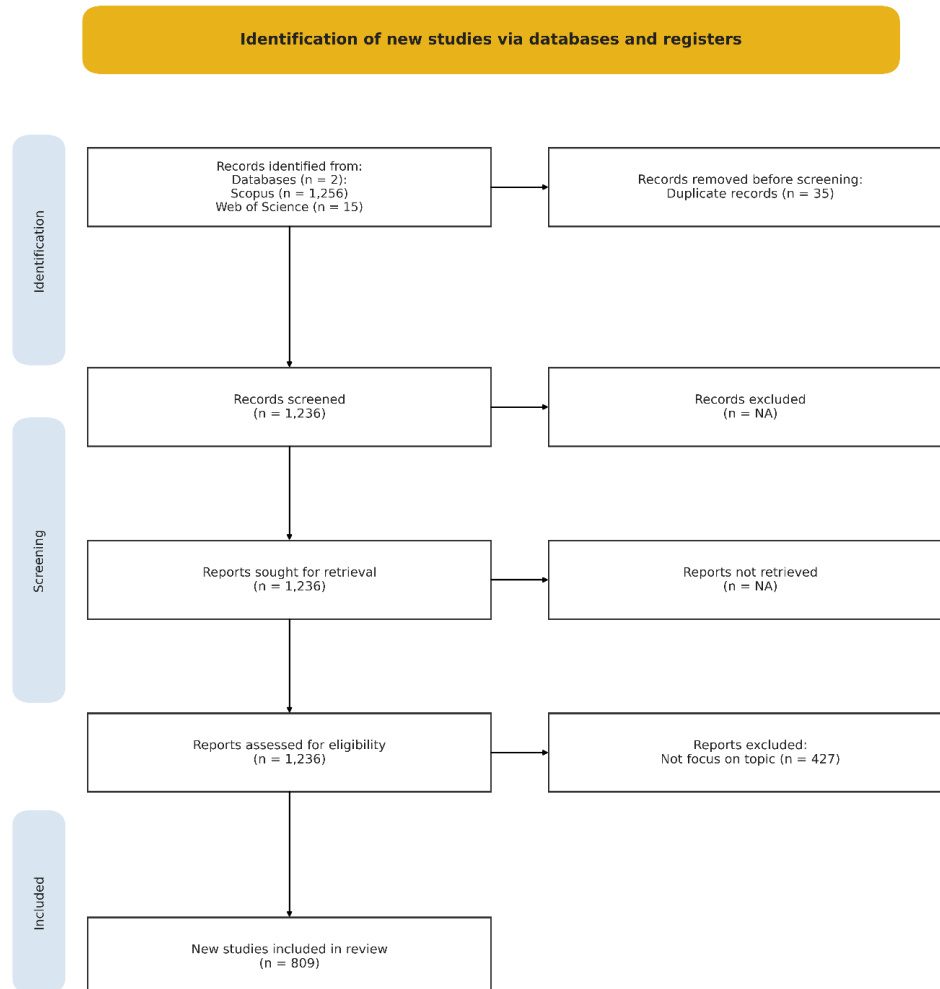
contexts or lacked substantive discussion of governance, risk, or organizational implications. Following this process, a final sample of 809 articles was retained for analysis. These articles formed the empirical basis of the Systematic Literature Review.

For each selected study, structured data extraction was conducted to capture relevant bibliographic and analytical information, including authorship, publication year, journal, title, abstract,

DOI, and country of origin. This information supported both descriptive analysis and the subsequent thematic classification aligned with the research questions. Fig. 2 presents the flow of the selection process following the PRISMA framework, ensuring transparency and replicability of the review procedure (Page et al., 2021).

A substantial number of articles were excluded during this phase because their orientation was not strictly aligned with the managerial and organizational focus of

Figure 2. PRISMA Flow Diagram



the study. In particular, studies emphasizing algorithmic performance, system architecture, or technical experimentation without explicit linkage to managerial decision-making, governance, or responsibility did not meet the quality and relevance criteria. This filtering process was essential to preserve the analytical integrity of the review and to ensure consistency with the theoretical framework.

Phase 5: Data Synthesis

The data synthesis process was designed to systematically address the five research questions (RQ1–RQ5) guiding the review. A qualitative content analysis approach was employed to draw structured and objective inferences from the information reported in the selected articles, following established SLR methodologies in software engineering and management research.

The synthesis was conducted in two complementary phases. First, a mechanical coding phase was used to classify each article into predefined analytical categories corresponding to the research questions, such as primary managerial purpose of AI use, type of organizational AI governance, human–AI decision reliance management, evidence of organizational value, and prominent managerial risk domains. Second, an interpretive analysis phase was undertaken to examine how effectively the content of each article addressed the research questions and to identify recurring patterns, tensions, and gaps across the literature.

To support analytical clarity and comparability, graphical representations were developed to summarize the

distribution of categories across the corpus, facilitating cross-RQ analysis and visual synthesis of results. In addition, article objectives, keywords, and conceptual framings were examined to identify convergences and emerging themes relevant to managerial practice and theory development. This structured synthesis provided the empirical foundation for the results, discussion, and conclusions presented in the subsequent sections.

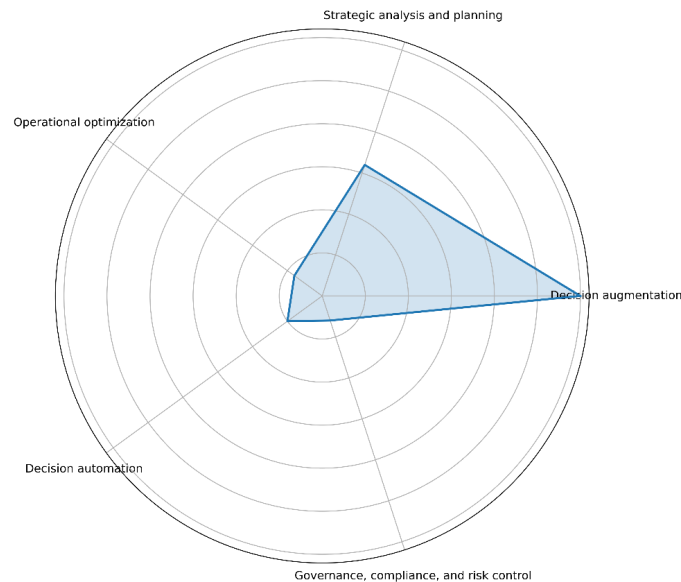
4. Results

4.1 Primary Managerial Purpose of AI Use

Fig. 3 presents the distribution of the primary managerial purposes associated with the use of artificial intelligence in the analyzed literature, synthesizing how the reviewed studies conceptualize the role of AI within organizational decision-making processes. This visualization allows for the identification of dominant patterns regarding whether AI is used as a support tool, a partial substitute, or a control mechanism within business management, providing a first empirical approximation to the predominant functional approach in recent research (Akiba & Fraboni, 2023).

The analysis of the graph shows a clear predominance of decision augmentation, which indicates that most studies conceive AI as an instrument of cognitive support for managers rather than as an autonomous substitute for human judgment. This trend is consistent with studies that underscore the value of AI in improving informational quality, reducing bias, and expanding analytical capacity without completely displacing managerial responsibility (Abramski et al., 2023). Likewise, the

Figure 3. Primary Managerial Purpose of AI Use

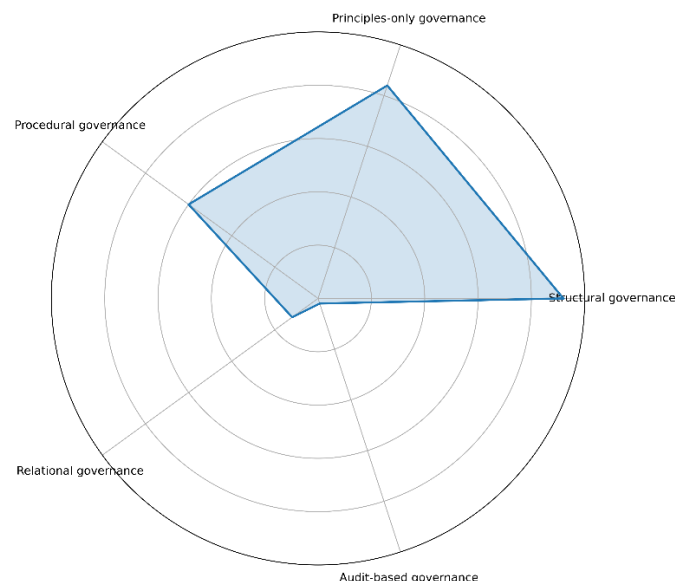


relevant presence of categories linked to strategic planning suggests that AI is increasingly associated with medium- and long-term decisions, and not only with operational tasks.

In contrast, the categories related to decision automation and regulatory control

have a lower relative representation, which suggests a persistent caution in the full delegation of critical decisions to algorithmic systems (Wang, 2021). This finding reinforces the idea that, from a managerial perspective, the adoption of AI continues to be predominantly hybrid, keeping the human decision-maker as a

Figure 4. Type of Organizational AI Governance



central actor and opening the way to the analysis of how this interaction is managed, an issue that is further examined in the following RQs.

4.2 Type of Organizational AI Governance

Fig. 4 shows the distribution of the types of organizational AI governance identified in the literature, allowing for the observation of how organizations structure, regulate, and supervise the use of algorithmic systems in business contexts. This figure is key to understanding whether governance is conceived as a formal set of structures, as operating procedures, or as general principles of action (Henderson et al., 2019).

The results reflect a strong concentration on procedural governance models, which indicates that most studies emphasize internal policies, protocols of use, and operational guidelines as the main control mechanisms. This pattern is consistent with studies that highlight the need to integrate AI into existing risk management

and internal control systems, rather than to create entirely new structures (Himabindu et al., 2023). Structural governance, although present, appears as complementary and not as the dominant axis.

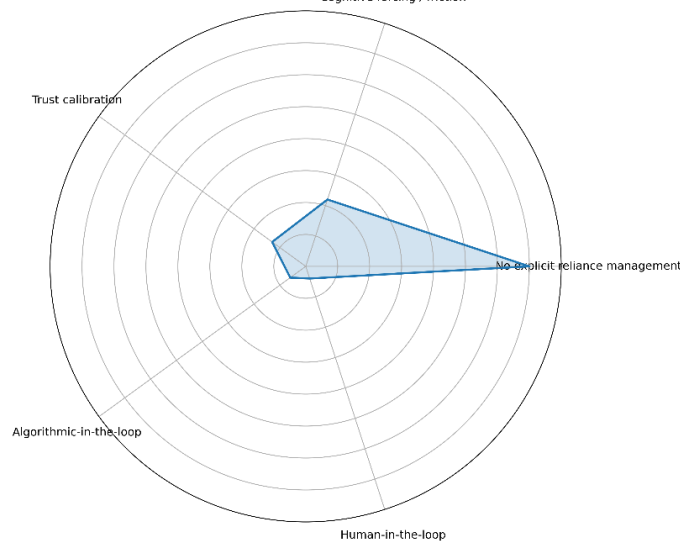
On the other hand, approaches based exclusively on ethical principles or abstract normative statements show a lower presence, suggesting a transition from declarative frameworks to more operational and auditable mechanisms (Landers & Behrend, 2023). This finding points to a maturation of the managerial discourse on AI, in which governance ceases to be merely aspirational and becomes a functional component of management, thereby connecting directly with the implementation challenges analyzed in the following category.

4.3 Human-AI Decision Reliance Management

Fig. 5 presents how the literature addresses

Figure 5. Human-AI Decision Reliance Management

Fig. 5 Human-AI Decision Reliance Management
Cognitive forcing / friction

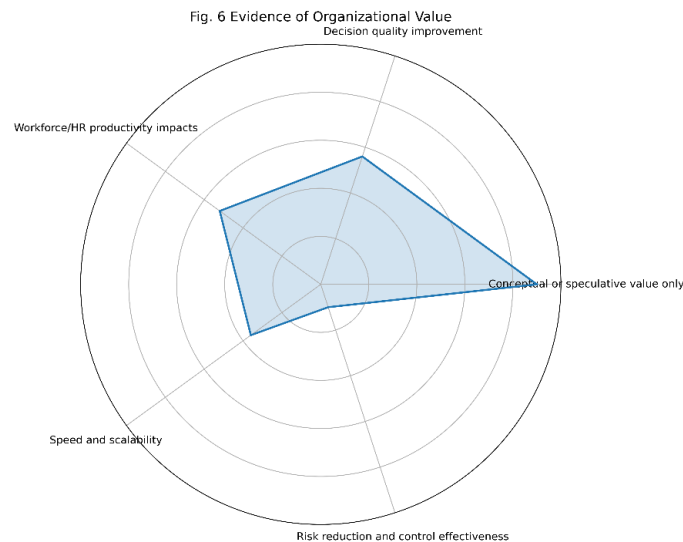


the management of reliance between humans and AI systems in decision-making, a central aspect to understand the risks of overtrust, underutilization, or inappropriate delegation in organizational contexts (Golovianko et al., 2023). This figure allows for the identification of whether the studies explicitly consider mechanisms to calibrate trust and the role of the human decision-maker.

The dominant pattern shows a high presence of human-in-the-loop approaches,

However, it is notable that a significant proportion of articles do not explicitly describe clear mechanisms for managing reliance, suggesting a gap between the technical development of AI and its deliberate integration into management processes (Alharthi, 2025). This absence reinforces the need for conceptual frameworks that more systematically articulate human-AI interaction, an issue that is directly connected to the evaluation of organizational value presented in Figure 6.

Figure 6. Evidence of Organizational Value



which confirms that most studies advocate maintaining the human manager as an active supervisor of the decision-making process. This preference is aligned with research that warns of the risks of excessive automation and underscores the need to preserve human agency to ensure accountability and organizational sense (García-López et al., 2025). Likewise, attention to the calibration of trust emerges as a recurring theme, especially in studies that analyze the adoption of AI from the perspective of organizational behavior.

4.4 Evidence of Organizational Value

Fig. 6 summarizes the types of organizational value that the literature attributes to the use of AI in managerial decision-making, allowing a distinction to be made between empirically measured benefits and contributions of a conceptual or speculative nature (Ahmad et al., 2023). This visualization is critical to assessing the extent to which the promises of AI translate into concrete organizational outcomes.

The analysis reveals that improved decision quality and operational efficiency are the most frequently reported benefits, which coincides with empirical studies documenting increases in decision-making accuracy, speed, and consistency (Brajovic et al., 2023). These results reinforce the narrative of AI as a tool to optimize management processes in complex and dynamic environments.

However, a relevant part of the literature is limited to discussing the potential value of AI without providing direct empirical evidence, suggesting that the field is still in a consolidation phase (Chan & Zhou, 2023). This coexistence of measured results and theoretical promises highlights the importance of analyzing the associated risks and limitations, an aspect that is explicitly addressed in Figure 7.

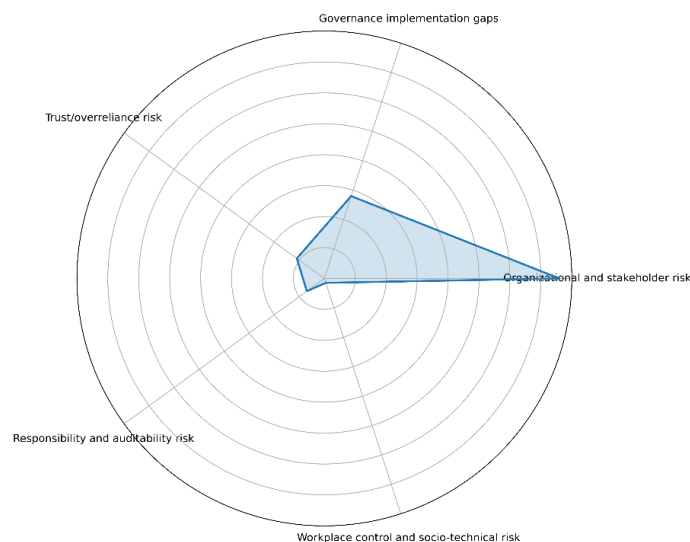
4.5 Foregrounded Managerial Risk Domain
Figure 7 presents the most emphasized managerial risk domains in the AI and decision-making literature, providing

an integrated view of the concerns that accompany their adoption in business contexts (Liang et al., 2022). This figure allows for the identification of whether the risks are conceptualized mainly in technical, organizational or liability terms.

The results show that risks related to responsibility and auditability are at the heart of the matter, reflecting a widespread concern about the traceability of algorithmic decisions and the allocation of responsibilities (Pushkarna et al., 2022). This emphasis is consistent with work that points out the difficulty of auditing complex systems and the need to maintain clear accountability mechanisms in automated environments.

Likewise, the risks associated with excessive trust and poor organizational implementation appear recurrently, underlining that the challenges of AI are not only technological but deeply managerial and socio-organizational (Méndez Carpio et al., 2023). Overall, this figure reinforces

Figure 7. Foregrounded Managerial Risk Domain



the idea that the adoption of AI in business administration requires not only technical capabilities but also robust management frameworks, an issue that will be taken up in an integrative way in the discussion section.

5. Discussion

5.1 Primary Managerial Purpose of AI Use

The results associated with RQ1 confirm that the use of artificial intelligence in the management literature is predominantly aligned with the construct of Augmented Managerial decision-making, reinforcing the idea that AI is mostly conceived as a mechanism to support human judgment (Akiba & Fraboni, 2023). The high concentration of decision augmentation studies coincides with theoretical approaches that highlight the complementarity between human and algorithmic capabilities, particularly in environments characterized by complexity and uncertainty (Duan et al., 2023).

This finding contrasts with more techno-deterministic discourses that promote the full automation of decisions, which are underrepresented in the corpus analyzed (Golovianko et al., 2023). The lower presence of decision automation approaches suggests that, from a managerial perspective, there is still structural caution against the full delegation of strategic decisions to algorithmic systems (Ananny & Crawford, 2018). This caution has been documented in studies that point to risks of loss of control, erosion of responsibility, and difficulties in justifying automated decisions to multiple stakeholders.

Overall, the results of RQ1 empirically support the first construct and show that the current literature privileges hybrid decision models (Brajovic et al., 2023). This emphasis lays the foundation for analyzing how organizations attempt to govern and structure this increased use of AI, an issue that is directly addressed in RQ2.

5.2 Type of Organizational AI Governance

The distribution observed in RQ2 shows a clear convergence between the empirical results and the construct of Organizational Governance of AI in Managerial Contexts, particularly in relation to the predominance of procedural governance approaches (Silva, 2022). The literature reviewed emphasizes internal policies, protocols, and operational mechanisms as the main instruments to regulate the use of AI, which coincides with research that points to the need to integrate algorithmic governance into traditional systems of organizational control (Henderson et al., 2019).

This procedural emphasis contrasts with approaches based exclusively on general ethical principles, which appear with less relative weight in the corpus (Black et al., 2022). This result suggests an evolution from abstract regulatory frameworks to more pragmatic and applicable solutions aimed at the effective implementation of AI in real business contexts (Abunaser et al., 2025; Ahmad et al., 2023). Recent studies have pointed out that, without clear procedures, ethical principles lack operational capacity and fail to substantially influence managerial decision-making processes.

Thus, the results of RQ2 reinforce the second construct by showing that AI governance is conceived primarily as an organizational problem and not merely as a normative one (Landers & Behrend, 2023). This procedural approach creates the conditions for analyzing how the risks and responsibilities associated with human-AI interaction are managed in practice, which is explored in the following RQs.

5.3 Human-AI Decision Reliance Management

The findings of RQ3 show a strong alignment with the construct of managerial risk and responsibility in human-AI decision systems, particularly with regard to the maintenance of the human decision-maker as an active supervisor (Ahmed Ali Linkon et al., 2024). The predominance of human-in-the-loop approaches suggests that the literature explicitly recognizes the risks of over-reliance and advocates preserving human agency in critical decision-making processes (Méndez Carpio et al., 2023).

However, it is significant that a considerable proportion of articles do not explicitly describe clear mechanisms for managing human-AI decision reliance (Ben-Michael et al., 2024). This absence highlights a gap between the theoretical recognition of risks and their operationalization in concrete management practices. Previous research has warned that the lack of clear trust-calibration strategies can lead to both overconfidence and underutilization of algorithmic systems, negatively affecting decision-making quality (Duan et al., 2023).

In this sense, the results of RQ3 not only confirm the third construct but also reveal

an opportunity for theoretical and practical development. The explicit management of human-AI reliance emerges as a critical area that requires greater attention in future research, especially in contexts where AI is integrated into strategic decisions with high organizational impact (Chang et al., 2023).

5.4 Evidence of Organizational Value

The results associated with RQ4 reinforce the construct of augmented managerial decision-making by showing that the literature attributes the organizational value of AI mainly to improvements in decision quality and operational efficiency (Chan & Zhou, 2023). These findings are consistent with empirical studies documenting increases in accuracy, speed, and decision consistency when AI is used as analytical support for managers (Wang, 2021).

However, the relevant presence of works that discuss the value of AI conceptually or speculatively suggests that the field is still in a consolidation phase (Jiang et al., 2024). This coexistence of empirical evidence and theoretical promise has been pointed out by various authors as a typical characteristic of emerging technologies, where expectations tend to precede the systematic measurement of results (Batool et al., 2024).

Thus, RQ4 shows that, although the managerial value of AI is widely recognized, the need for more robust studies that directly link the use of AI to measurable organizational outcomes persists (Ahmad et al., 2023). This empirical gap is connected to the risks and limitations discussed in

RQ5, thereby closing the analytical cycle of the study.

5.5 Foregrounded Managerial Risk Domain

The results of RQ5 strongly confirm the third construct, showing that responsibility and auditability risks are at the heart of the management literature on AI (Henderson et al., 2019). Concerns about the traceability of algorithmic decisions and the allocation of responsibilities reflect a growing awareness that AI adoption redefines traditional structures of control and responsibility (Lian et al., 2022).

Likewise, the relevance of the risks associated with excessive trust and poor organizational implementation underscores that the challenges of AI are profoundly socio-organizational (Ahmad et al., 2023; Díaz-Rodríguez et al., 2023). Recent studies highlight that failures in AI adoption are rarely due exclusively to technical limitations, but rather to issues of organizational alignment, culture, and managerial capabilities (Abunaser et al., 2025).

Taken together, RQ5 closes the discussion by showing that risk management is not a peripheral aspect but a central component of the managerial use of AI (Alharthi, 2025). This result reinforces the need for integrated approaches that combine decision-making augmentation, organizational governance, and managerial responsibility, providing a solid basis for the study's conclusions.

6. Conclusions

The objective of this study was to

systematically analyze how recent literature addresses the use of artificial intelligence in managerial decision-making, integrating the constructs of augmented managerial decision-making, organizational AI governance and managerial risk and responsibility. The results show that contemporary research converges on a hybrid conception of AI, in which algorithmic systems are mainly understood as tools to support human judgment and not as autonomous substitutes for the decision-maker. This finding confirms and reinforces theoretical approaches that maintain that the strategic value of AI emerges when it is integrated into existing decision-making processes, respecting the structures of responsibility and control typical of business administration.

Likewise, the analysis of the five research questions shows that the organizational governance of AI has progressively shifted towards procedural and operational approaches, oriented towards practical implementation rather than the formulation of abstract principles. This trend suggests a maturation of the field, in which organizations seek to translate ethical and strategic values into concrete mechanisms of supervision, control, and responsibility. However, the results also reveal that such governance is not always accompanied by explicit strategies to manage human-AI decision reliance, which generates tensions between decision-making efficiency and managerial responsibility.

Overall, this study makes a distinctive contribution by offering an integrated managerial reading of the literature on artificial intelligence, overcoming

fragmented approaches that analyze technology, ethics, or strategy in isolation. By articulating the three proposed constructs with systematically classified empirical evidence, the work provides an analytical framework that allows for an understanding not only how AI is used in business decision-making but also under what organizational conditions such use generates value and what risks must be managed. This integration strengthens the relevance of the study for both academic research and managerial practice.

However, the study has some limitations that must be recognized. First, the analysis relies on secondary information extracted from titles, abstracts, and keywords, which can hide nuances present in the full text of the articles. Second, although the corpus analyzed is large and diverse, the literature reviewed presents considerable heterogeneity in terms of methodologies and organizational contexts, which makes direct comparisons of findings difficult. Finally, the classification of articles, although based on systematic and replicable rules, necessarily implies a degree of interpretation that could be refined through more in-depth qualitative analyses.

These limitations open several avenues for future research. First, a greater number of empirical studies are required that explicitly link the use of AI with measurable organizational outcomes, particularly in high-impact strategic decisions. Second, future research could delve into the specific mechanisms of managing human-AI reliance, exploring how trust, oversight, and responsibility are calibrated in different

sectoral contexts. Finally, it is pertinent to move towards integrated models that combine governance, organizational design, and managerial capabilities, in order to develop more robust frameworks for the responsible and effective adoption of artificial intelligence in business administration.

Sources of information

- Abhishek, A., Erickson, L., & Bandopadhyay, T. (2025). Data and AI governance: Promoting equity, ethics, and fairness in large language models (Versión 1). arXiv. <https://doi.org/10.48550/ARXIV.2508.03970>
- Abramski, K., Citraro, S., Lombardi, L., Rossetti, G., & Stella, M. (2023). Cognitive Network Science Reveals Bias in GPT-3, GPT-3.5 Turbo, and GPT-4 Mirroring Math Anxiety in High-School Students. *Big Data and Cognitive Computing*, 7(3), 124. <https://doi.org/10.3390/bdcc7030124>
- Abunaser, F. M., Hamd, M. M. M., Bani-Oraba, A. M. N., Hamed, O., Alshiyab, M. Q. M., & Shebani, Z. (2025). Dynamic Capabilities of University Administration and Their Impact on Student Awareness of Artificial Intelligence Tools. *Sustainability*, 17(15), 7092. <https://doi.org/10.3390/su17157092>
- Acosta-Enriquez, B. G., Ballesteros, M. A. A., De Los Angeles Guzman Valle, M., Angaspilco, J. E. M., Blanco-García, L. E., Ventura, G. C., Requejo, J. D. C., & Torre, M. C. (2025). Determinants of AI Use in University Teachers: The Role of Leadership, Teaching Concerns, and Constructivist Pedagogical Beliefs. *Human Behavior and Emerging Technologies*, 2025(1), 4834893. <https://doi.org/10.1155/hbe2/4834893>
- Afroogh, S., Akbari, A., Malone, E., Kargar, M., & Alambeigi, H. (2024). Trust in AI: Progress, Challenges, and Future Directions (Versión 3). arXiv. <https://doi.org/10.48550/ARXIV.2403.14680>
- Ahmad, M., Alhalaiqa, F., & Subih, M. (2023). Constructing and testing the psychometrics of an instrument to measure the attitudes, benefits, and threats associated with the use of Artificial Intelligence tools in higher education. *Journal of Applied Learning and Teaching*, 6(2), 114-120. Scopus. <https://doi.org/10.37074/jalt.2023.6.2.36>
- Ahmed Ali Linkon, Mujiba Shaima, Md Shohail Uddin Sarker, Badruddowza, Norun Nabi, Md Nasir Uddin Rana, Sandip Kumar Ghosh, Hammed Esa, & Faiaz Rahat Chowdhury. (2024). Advancements and Applications of Generative Artificial Intelligence and Large Language Models on Business Management: A Comprehensive Review. *Journal of Computer Science and Technology Studies*, 6(1), 225-232. <https://doi.org/10.32996/jcsts.2024.6.1.26>
- Akiba, D., & Fraboni, M. C. (2023). AI-Supported Academic Advising: Exploring ChatGPT's Current State and Future Potential toward Student Empowerment. *Education Sciences*, 13(9). Scopus. <https://doi.org/10.3390/educsci13090885>
- Alharthi, S. (2025). Harnessing Knowledge: The Robust Role of Knowledge Management Practices and Business Intelligence Systems in Developing Entrepreneurial Leadership and Organizational Sustainability in SMEs. *Sustainability*, 17(14), 6264. <https://doi.org/10.3390/su17146264>

- Al-Kfairy, M., Mustafa, D. G., Kshetri, N., Insiew, M., & Alfandi, O. (2024). Ethical challenges and solutions of generative AI: An interdisciplinary perspective. *Informatics*. <https://doi.org/10.3390/informatics11030058>
- Ameen, N., Tarhini, A., Reppel, A., & Anand, A. (2021). Customer experiences in the age of artificial intelligence. *Computers in Human Behavior*, 114, 106548. <https://doi.org/10.1016/j.chb.2020.106548>
- Ananny, M., & Crawford, K. (2018). Seeing without knowing: Limitations of the transparency ideal and its application to algorithmic accountability. *New Media & Society*, 20(3), 973-989. <https://doi.org/10.1177/1461444816676645>
- Bacon, D. R. (2024). Recommendations for the Use of Experimental Designs in Management Education Research. *Journal of Management Education*, 48(6), 1121-1147. <https://doi.org/10.1177/10525629241252802>
- Baothman, F. A. (2021). An Intelligent Big Data Management System Using Haar Algorithm-Based Nao Agent Multisensory Communication. *Wireless Communications and Mobile Computing*, 2021. Scopus. <https://doi.org/10.1155/2021/9977751>
- Batool, A., Zowghi, D., & Bano, M. (2024). Responsible AI Governance: A Systematic Literature Review (Versión 1). arXiv. <https://doi.org/10.48550/ARXIV.2401.10896>
- Ben-Michael, E., Greiner, D. J., Huang, M., Imai, K., Jiang, Z., & Shin, S. (2024). Does AI help humans make better decisions? A statistical evaluation framework for experimental and observational studies (Versión 3). arXiv. <https://doi.org/10.48550/ARXIV.2403.12108>
- Black, S., Biderman, S., Hallahan, E., Anthony, Q., Gao, L., Golding, L., He, H., Leahy, C., McDonell, K., Phang, J., Pieler, M., Prashanth, U. S., Purohit, S., Reynolds, L., Tow, J., Wang, B., & Weinbach, S. (2022). GPT-NeoX-20B: An Open-Source Autoregressive Language Model. *Proceedings of BigScience Episode #5 -- Workshop on Challenges & Perspectives in Creating Large Language Models*, 95-136. <https://doi.org/10.18653/v1/2022.bigscience-1.9>
- Bommasani, R., Hudson, D. A., Adeli, E., Altman, R., Arora, S., von Arx, S., Bernstein, M. S., Bohg, J., Bosselut, A., Brunskill, E., Brynjolfsson, E., Buch, S., Card, D., Castellon, R., Chatterji, N., Chen, A., Creel, K., Davis, J. Q., Demszky, D., ... Liang, P. (2021). On the Opportunities and Risks of Foundation Models (Versión 3). arXiv. <https://doi.org/10.48550/ARXIV.2108.07258>
- Brajovic, D., Renner, N., Goebels, V. P., Wagner, P., Fresz, B., Biller, M., Klaeb, M., Kutz, J., Neuhuettler, J., & Huber, M. F. (2023). Model Reporting for Certifiable AI: A Proposal from Merging EU Regulation into AI Development (Versión 1). arXiv. <https://doi.org/10.48550/ARXIV.2307.11525>
- Buijsman, S., Carter, S. E., & Bermúdez, J.-P. (2025). Autonomy by Design: Preserving Human Autonomy in AI Decision-Support. *Philosophy &*

- Technology, 38(3), 97. <https://doi.org/10.1007/s13347-025-00932-2>
- Casey, B., Farhangi, A., & Vogl, R. (2019). Rethinking explainable machines: The gdpr's 'right to explanation' debate and the rise of algorithmic audits in enterprise. *Berkeley Technology Law Journal*, 34(1), 143. <https://doi.org/10.15779/Z38M32N986>
- Chan, C. K. Y., & Zhou, W. (2023). An expectancy value theory (EVT) based instrument for measuring student perceptions of generative AI. *Smart Learning Environments*, 10(1). Scopus. <https://doi.org/10.1186/s40561-023-00284-4>
- Chang, S.-H., Yao, K.-C., Chung, C.-Y., Nien, S.-C., Chen, Y.-T., Ho, W.-S., Lin, T.-C., Shih, F.-C., & Chung, T.-C. (2023). Integration of Artificial Intelligence and Machine Learning Content in Technology and Science Curriculum. *International Journal of Engineering Education*, 39(6), 1343-1357. Scopus.
- Cheng, K., & Wu, H. (2024). Policy framework for the utilization of generative AI. *Critical Care*, 28(1), 128.
- Del Pozo, C., Nuno Gomes de Andrade, N., & Rojas Arroyo, D. (2023). Prototipo de políticas públicas sobre transparencia y explicabilidad de sistemas de inteligencia artificial [Informe técnico]. OpenLoop. <https://openloop.org/reports/2023/10/Prototipo-de-Políticas-Públicas-sobre-Transparencia-y-Explicabilidad-de-Sistemas-de-IA.pdf>
- Díaz-Rodríguez, N., Del Ser, J., Coeckelbergh, M., de Prado, M. L., Herrera-Viedma, E., & Herrera, F. (2023). Connecting the Dots in Trustworthy Artificial Intelligence: From AI Principles, Ethics, and Key Requirements to Responsible AI Systems and Regulation (Versión 2). arXiv. <https://doi.org/10.48550/ARXIV.2305.02231>
- Ding, Z., Jiang, S., Xu, X., & Han, Y. (2022). An Internet of Things based scalable framework for disaster data management. *Journal of Safety Science and Resilience*, 3(2), 136-152. Scopus. <https://doi.org/10.1016/j.jnlssr.2021.10.005>
- Duan, W., Liu, Z., Jia, C., Wang, S., Ma, S., & Gao, W. (2023). Differential Weight Quantization for Multi-Model Compression. *IEEE Transactions on Multimedia*, 25, 6397-6410. <https://doi.org/10.1109/TMM.2022.3208530>
- European Union. (2024). Article 9: Risk management system. <https://artificialintelligenceact.eu/article/9/>
- for Economic Co-operation, O., & (OECD), D. (2025). AI openness: A primer for policymakers (No. 44; OECD artificial intelligence papers). OECD Publishing. <https://doi.org/10.1787/958d292b-en>
- García-López, I. M., Ramírez-Montoya, M. S., & Molina-Espinosa, J. M. (2025). Generative artificial intelligence in education: A systematic analysis of opportunities, challenges, and responses. *Interactive Learning Environments*, 1-24. <https://doi.org/10.1080/10494820.2025.2519133>
- Golovianko, M., Gryshko, S., Terziyan, V., & Tuunanen, T. (2023). Responsible

- cognitive digital clones as decision-makers: a design science research study. *European Journal of Information Systems*, 32(5), Article 5. Scopus. <https://bit.ly/3QFqaLY>
- Grimmelikhuijsen, S., & Meijer, A. (2022). Legitimacy of Algorithmic Decision-Making: Six Threats and the Need for a Calibrated Institutional Response. *Perspectives on Public Management and Governance*, 5(3), 232-242. <https://doi.org/10.1093/ppmgov/gvac008>
- Henderson, L. H., Wersun, A., Wilson, J., Mo-ching Yeung, S., & Zhang, K. (2019). Principles for responsible management education in 2068. *Futures*, 111, 81-89. Scopus. <https://doi.org/10.1016/j.futures.2019.05.005>
- Himabindu, M., V, R., Gupta, M., Rana, A., Chandra, P. K., & Abdulaali, H. S. (2023). Neuro-Symbolic AI: Integrating Symbolic Reasoning with Deep Learning. 2023 10th IEEE Uttar Pradesh Section International Conference on Electrical, Electronics and Computer Engineering (UPCON), 1587-1592. <https://doi.org/10.1109/UPCON59197.2023.10434380>
- Jiang, H., Zhang, X., Mahari, R., Kessler, D., Ma, E., August, T., Li, I., Pentland, A. «Sandy», Kim, Y., Roy, D., & Kabbara, J. (2024). Leveraging Large Language Models for Learning Complex Legal Concepts through Storytelling (Versión 4). arXiv. <https://doi.org/10.48550/ARXIV.2402.17019>
- Landers, R. N., & Behrend, T. S. (2023). Auditing the AI auditors: A framework for evaluating fairness and bias in high stakes AI predictive models. *American Psychologist*, 78(1), 36-49. <https://doi.org/10.1037/amp0000972>
- Larsson, S. (2020). On the Governance of Artificial Intelligence through Ethics Guidelines. *Asian Journal of Law and Society*, 7(3), 437-451. <https://doi.org/10.1017/als.2020.19>
- Li, Z., Zhu, H., Lu, Z., Xiao, Z., & Yin, M. (2025). From Text to Trust: Empowering AI-assisted Decision Making with Adaptive LLM-powered Analysis. *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, 1-18. <https://doi.org/10.1145/3706598.3713133>
- Liang, P., Bommasani, R., Lee, T., Tsipras, D., Soylu, D., Yasunaga, M., Zhang, Y., Narayanan, D., Wu, Y., Kumar, A., Newman, B., Yuan, B., Yan, B., Zhang, C., Cosgrove, C., Manning, C. D., Ré, C., Acosta-Navas, D., Hudson, D. A., ... Koreeda, Y. (2022). Holistic Evaluation of Language Models. <https://doi.org/10.48550/ARXIV.2211.09110>
- Méndez Carpio, C. R., Arévalo Medranda, S. R., León Segovia, L. S., Parra Guerrero, E. F., & Siguencia Tello, S. P. (2023). Tecnopatías y dependencias: Uso incorrecto de la tecnología. *Killkana Social*, 7(3), 77-88. <https://doi.org/10.26871/killkanasocial.v7i3.1407>
- Mišić, J., Van Est, R., & Kool, L. (2025). Good governance of public sector AI: A combined value framework for good order and a good society. *AI and Ethics*, 5(5), 4875-4889. <https://doi.org/10.1007/s43681-025-00751-3>

- Page, M. J., McKenzie, J. E., Bossuyt, P. M., Boutron, I., Hoffmann, T. C., Mulrow, C. D., Shamseer, L., Tetzlaff, J. M., Akl, E. A., Brennan, S. E., Chou, R., Glanville, J., Grimshaw, J. M., Hróbjartsson, A., Lalu, M. M., Li, T., Loder, E. W., Mayo-Wilson, E., McDonald, S., ... Moher, D. (2021). Declaración PRISMA 2020: Una guía actualizada para la publicación de revisiones sistemáticas. *Revista Española de Cardiología*, 74(9), Article 9. <https://doi.org/10.1016/j.recresp.2021.06.016>
- Paul, J., Lim, W. M., O'Cass, A., Hao, A. W., & Bresciani, S. (2021). Scientific procedures and rationales for systematic literature reviews (SPAR-4-SLR). *International Journal of Consumer Studies*, 45(4). <https://doi.org/10.1111/ijcs.12695>
- Pushkarna, M., Zaldivar, A., & Kjartansson, O. (2022). Data Cards: Purposeful and Transparent Dataset Documentation for Responsible AI (Versión 1). arXiv. <https://doi.org/10.48550/ARXIV.2204.01075>
- Sauer, P. C., & Seuring, S. (2023). How to conduct systematic literature reviews in management research: A guide in 6 steps and 14 decisions. *Review of Managerial Science*, 17(5), 1899-1933. <https://doi.org/10.1007/s11846-023-00668-3>
- Silva, A. A. (2022). Gobernanza, poder y autonomía universitaria en la era de la innovación. *Perfiles Educativos*, 44(178), Article 178. Scopus. <https://doi.org/10.22201/issue.24486167e.2022.178.60735>
- Wang, Y. (2021). When artificial intelligence meets educational leaders' data-informed decision-making: A cautionary tale. *Studies in Educational Evaluation*, 69, 100872. <https://doi.org/10.1016/j.stueduc.2020.100872>
- Zhai, C., Wibowo, S., & Li, L. D. (2024). The effects of over-reliance on AI dialogue systems on students' cognitive abilities: A systematic review. *Smart Learning Environments*, 11(1), 28. <https://doi.org/10.1186/s40561-024-00316-7>