

Generative Artificial Intelligence in Business Decision-Making: Emerging Frameworks, Enterprise Applications, and Future Challenges

M.SANGEETHA¹, Una Suman Kumar Patro², K.ARPITHA³, Piyal Roy⁴, Burri Naresh⁵, P. Mayavel⁶

¹Department Of COMMERCE, Chennai, Tamil Nadu, Sangeethamohanraj.m@gmail.com

²Senior Data Engineer, samsumankumarpatro@gmail.com

³CSE (AI and ML), Geethanjali College of Engineering and Technology, Keesara, Cheeryal, Telangana, drarpitha.cse@gcet.edu.in

⁴CSE-AI, Brainware University, 24 PGS(N), Kolkata, West Bengal, piyalroy000@gmail.com

⁵CSE(Cyber Security), CVR College of Engineering, naresh.burri@gmail.com

⁶Department of Mathematics, Vel Tech Rangarajan Dr. Sagunthala R&D Institute of Science and Technology, Avadi, Chennai, India, mayavel.maya@gmail.com

Abstract: Generative artificial intelligence (GenAI) has moved from experimental novelty to a central input in organizational decision-making, with McKinsey's Q1 2026 Global AI Survey finding that 65 percent of organizations now use generative AI in at least one business function, roughly double the adoption rate reported ten months earlier, and 72 percent report at least one AI workload in production. Despite this scale of adoption, the empirical record on business value remains sharply divided. This paper synthesizes recent academic and industry evidence, drawing on 24 sources published primarily between 2023 and 2026, to examine three interlocking questions: what theoretical frameworks currently explain GenAI's role in managerial and strategic decision-making, how enterprises are applying GenAI in practice across functional areas, and what structural challenges limit the translation of GenAI adoption into measurable business value. The paper synthesizes evidence from strategic management research on AI-assisted evaluation of business alternatives, organizational theory on GenAI's emerging roles in decision processes, and empirical field studies on ambiguity handling and sycophantic behavior in AI-generated business advice, alongside a widely cited 2025 MIT study finding that 95 percent of enterprise generative AI pilots fail to deliver measurable profit-and-loss impact. Findings indicate that GenAI functions most reliably as an augmentation tool that aggregates and structures diverse inputs for human judgment, rather than as an autonomous decision-maker, that single-model evaluations of strategic alternatives are frequently inconsistent and biased while aggregated multi-model evaluations approximate expert human judgment, and that the primary barrier to enterprise value is organizational and workflow integration rather than model capability. The paper concludes with a proposed decision-integration framework and implications for executives, AI governance functions, and researchers.

Keywords: generative artificial intelligence, business decision-making, enterprise AI, large language models, strategic management, AI governance, organizational adoption

1. Introduction

The release of ChatGPT in November 2022 marked a discontinuity in how generative artificial intelligence entered organizational life, moving the technology from research laboratories directly onto the desks of managers, analysts, and executives within a matter of months. Since then, enterprise adoption has accelerated sharply: 65 percent of organizations now report using generative AI in at least one business function, a figure that has roughly doubled in ten months according to McKinsey's Q1 2026 Global AI Survey, and 92 percent of Fortune 500 companies have adopted generative AI in some form. Global spending on generative AI is projected to reach 2.5 billion dollars in 2026, a fourfold increase over 2025 according to Gartner's forecasts, and Gartner separately projects that more than 60 percent of enterprise applications will embed generative AI to augment workflows by the end of 2026.



Yet this scale of adoption sits uneasily alongside a starkly different empirical picture of business value. A widely cited 2025 study from MIT's NANDA initiative, based on 52 executive interviews, a survey of 153 business leaders, and analysis of 300 public AI deployments, found that 95 percent of generative AI pilots deliver no measurable profit-and-loss impact, with only 5 percent of deployed systems creating significant, sustained value. This gap between adoption breadth and value depth, which the MIT researchers term the "GenAI Divide," is not primarily attributable to model quality or capability, but to a "learning gap" in how organizations integrate these tools into actual decision workflows. This tension between rapid, near-universal adoption and highly uneven realized value is the central empirical puzzle this paper addresses.

This paper pursues three interlocking research questions. First, what theoretical and conceptual frameworks have emerged to explain the role generative AI plays in managerial and strategic decision-making, and how does this role differ from earlier generations of decision-support and predictive-analytics tools. Second, how are enterprises applying generative AI in practice across core business functions, and what does the empirical evidence suggest about the reliability of AI-generated business recommendations. Third, what structural, technical, and organizational challenges currently limit the translation of generative AI adoption into measurable business value, and what does the emerging governance and risk-management literature suggest about addressing them. Answering these questions matters because organizations are currently making capital allocation decisions, Gartner-projected global spending in the billions, based on assumptions about generative AI's decision-making value that recent empirical evidence suggests are, for the overwhelming majority of current deployments, not being realized.

2. Literature Review

2.1 Theoretical Frameworks: From Automation to Augmentation

Early theoretical work on generative AI in organizational decision-making has converged on a shared distinction between AI as an automation substitute for human judgment and AI as an augmentation partner that expands the informational and cognitive resources available to human decision-makers. A recent conceptual synthesis positions generative AI as an agent of value co-creation in strategic processes, explicitly moving beyond its conventional framing as a passive decision-support instrument, and proposes testable hypotheses regarding GenAI's impact on strategic foresight, resource orchestration, and organizational resilience (Toha et al., 2026). This augmentation framing is echoed in a broader review of organizational decision-making theory, which finds that generative AI increasingly performs emerging roles beyond simple information retrieval, including balancing knowledge disparities between small and large firms, synthesizing structured and unstructured data into usable organizational knowledge, and exerting a democratizing effect on decision-making processes by increasing the availability of diverse information sources to decision-makers who previously lacked access to specialist analytical capacity (Rethinking organizational decision-making, 2026).

Classical decision-making theory continues to provide the interpretive scaffolding for this literature. One conceptual study synthesizes bounded rationality theory, Mintzberg's managerial roles framework, and socio-technical systems theory to analyze how AI reshapes human judgment, strategic foresight, and leadership dynamics, arguing that generative AI's core theoretical contribution is not replacing bounded rationality but reshaping the boundaries within which it operates, by expanding the volume and diversity of information a manager can process before a decision must be made (Generative Artificial Intelligence and Evaluating Strategic Decisions, 2024). This connects directly to Herbert Simon's foundational work on bounded rationality, which generative AI research increasingly treats as the natural theoretical anchor for understanding why AI-augmented managers might reach different, and not necessarily worse, decisions than either unaided humans or fully automated systems (Toha et al., 2026).

2.2 Empirical Evidence on the Quality of AI-Generated Strategic Judgments

A pivotal empirical contribution to this literature comes from a *Strategic Management Journal* study examining how well large language models evaluate strategic business alternatives compared with human experts. Using two studies, each analyzing sixty business models, one AI-generated and one drawn from a startup competition, the

researchers found that single AI evaluations of strategic alternatives are often inconsistent and exhibit systematic bias, but that aggregating evaluations across multiple large language models, assumed model roles, and prompt variations produces rankings that closely resemble those of human experts (Doshi et al., 2025). This is arguably the most methodologically rigorous evidence currently available on a question central to this paper: not whether generative AI can evaluate strategic decisions at all, but under what conditions its evaluations become reliable enough for managers to act on. The paper's practical conclusion, that managers should combine multiple AI evaluations rather than trust single-model outputs, provides one of the more directly actionable design principles in the reviewed literature.

A complementary and more cautionary empirical study examines generative AI's behavior specifically in ambiguous business decision contexts using a novel four-dimensional business ambiguity taxonomy and a human-in-the-loop experiment spanning strategic, tactical, and operational scenarios (Generative AI in Managerial Decision-Making, 2026). This study found that large language models frequently fail to recognize ambiguous inputs, instead resolving ambiguity silently and confidently rather than flagging it for human clarification, and additionally documented sycophantic behavior, a tendency to agree with and elaborate on flawed directives given by the user rather than challenging them, which the authors identify as a specific and previously underexamined failure mode in AI-assisted managerial decision-making (Generative AI in Managerial Decision-Making, 2026). This finding is theoretically important because it identifies a failure mode that a pure augmentation framework does not adequately anticipate: an AI system that agrees too readily with a flawed human premise may be worse than no AI assistance at all, since it can lend false confidence to a poor decision.

2.3 Entrepreneurial and Individual-Level Decision-Making

Beyond large-firm strategic contexts, generative AI's role in decision-making has also been studied at the level of individual and entrepreneurial decision-makers. A participant-observation study of latent entrepreneurs, individuals approaching entrepreneurship for the first time, interacting with generative AI to make venture-related decisions identified a central tension the authors describe as human-machine interaction along a spectrum of "taking or giving" control, developing a conceptual framework contributing to the innovation, decision-making, and entrepreneurship literature (Bastone et al., 2026). A separate integrative review of generative AI use in entrepreneurship proposes an "empowerment-entrapment" framework, arguing that the same properties that empower novice decision-makers with rapid access to sophisticated analysis, low cost, and broad domain coverage, can simultaneously entrap them in a dependency that atrophies independent judgment, a tension the review connects to broader psychological research on how individuals may use AI to avoid effortful cognitive engagement with difficult decisions (Generative AI Use in Entrepreneurship, 2026).

2.4 Enterprise Applications Across Functional Areas

The applied literature documents generative AI deployment across a widening range of business functions. Industry analysis characterizes 2026 enterprise adoption as having moved beyond early productivity-focused pilots toward programs that automate complex workflows, improve decision-making, and generate measurable operational efficiencies, with automation of report generation, document summarization, and routine administrative work cited as the most mature use cases (eClerx, 2026). Lucidworks' 2025 AI Benchmark Study, surveying more than 1,600 AI leaders and independently analyzing more than 1,100 company deployments, found that more than seven in ten organizations have introduced generative AI into operations, yet only 6 percent have fully implemented agentic AI, systems capable of autonomous multi-step task execution, indicating that most current enterprise deployment remains at the level of single-turn generation and retrieval rather than autonomous decision execution (Lucidworks, 2026). The same study found that 83 percent of AI leaders reported major or extreme concern about generative AI risk in 2025, an eightfold increase in just two years, with implementation cost overruns, data security, unreliable outputs, and lack of decision transparency cited as the dominant concerns (Lucidworks, 2026).

Financial services and sales and marketing functions have emerged as leading, though not necessarily most successful, adoption areas. Industry analysis citing Deloitte and McKinsey data finds financial services reporting a 4.2 times return on generative AI investment where implementation succeeds, yet notes that 80 percent of organizations

across sectors report no measurable earnings-before-interest-and-tax impact from their generative AI investments overall, and only 7 percent have achieved full scaling beyond pilot status (Automely, 2026). This pattern is corroborated by the MIT NANDA study, which found that AI budgets are disproportionately allocated to sales and marketing functions, despite operations and finance functions demonstrating better realized return on investment, a misallocation the report attributes to the higher visibility and executive appeal of customer-facing AI applications relative to less visible back-office automation (MIT NANDA, 2025).

2.5 The GenAI Divide: Structural Barriers to Enterprise Value

The MIT NANDA report's central contribution is identifying four structural factors that separate the small minority of successful generative AI deployments from the large majority that fail to scale (MIT NANDA, 2025). First, disruption remains limited to specific sectors: only two of nine major sectors studied, technology and media, show material business transformation attributable to generative AI use, while most other sectors report adoption without corresponding transformation. Second, an enterprise paradox exists in which large firms, defined in the study as those with more than 100 million dollars in annual revenue, lead in pilot volume and staffing allocated to AI initiatives, yet report the lowest rates of successful pilot-to-scale conversion, while mid-market firms move faster and more decisively, with top-performing organizations reporting average timelines of 90 days from pilot to full implementation. Third, an implementation-partnership advantage exists: tools built through partnerships blending internal AI specialists with external vendor expertise achieve a 67 percent success rate, compared with only 22 percent for internally built solutions, and externally built tools overall succeed at roughly twice the rate of internal builds. Fourth, and most fundamentally, the report identifies a "learning gap" as the core barrier to scaling, not infrastructure, regulation, or talent availability, but the inability of most deployed systems to retain feedback, adapt to organizational context, and improve over time, meaning that generic tools capable of high adoption for simple tasks systematically fail once workflows demand sustained contextual learning (MIT NANDA, 2025).

This structural diagnosis is reinforced by evidence of a parallel "shadow AI economy," in which employees at more than 90 percent of firms studied use personal, unofficial generative AI tools even when their employer's official pilots have stalled or failed, suggesting that individual-level utility from generative AI substantially outpaces organizational-level capacity to capture and formalize that utility into scaled decision-support systems (MIT NANDA, 2025).

2.6 Governance, Hallucination Risk, and Trust in High-Stakes Decisions

A distinct but closely related literature addresses the specific risks generative AI introduces when its outputs inform consequential business or public-sector decisions, centering on the problem of hallucination, the generation of fluent, confidently stated but factually unsupported content. A governance framework developed for large language models in enterprise analytics contexts argues that existing AI ethics and alignment principles, while valuable, have not been adequately translated into the specific metrics, controls, and workflows analytics teams need to design, deploy, and monitor LLM-powered decision-support systems in practice, and proposes systematically logging failure cases such as hallucinated regulatory citations or incorrect threshold logic into a model-level risk register used to refine acceptance criteria over time (From Models to Metrics, 2026). A related risk-oriented framework developed for public-sector large language model systems reframes hallucination management explicitly as a governance problem rather than a pure accuracy problem, arguing that institutions should not pursue the unrealistic goal of eliminating hallucination entirely, given current generative architectures, but should instead identify, measure, and calibrate institutional response to hallucination risk according to the consequential importance of the specific decision or procedure involved (Risk-Oriented Verification and Validation Framework, 2026).

Sector-specific governance architectures are beginning to formalize these principles into auditable structures. A prudential reliability framework developed for the reinsurance industry proposes a five-pillar architecture, governance, data lineage, assurance, resilience, and regulatory alignment, translating supervisory expectations from established financial regulatory regimes into measurable lifecycle controls for large language model deployment, implemented through a dedicated reliability and assurance benchmark evaluating whether governance-embedded

models meet prudential standards for grounding, transparency, and accountability (Dong, 2025). A broader bibliometric review of large language model hallucination research finds that effective organizational integration of generative AI requires both technical measures, including verification protocols and privacy protection, and human-centered organizational capabilities, including AI literacy training and responsible-use norms, echoing the MIT NANDA report's finding that the primary barrier to value realization is organizational rather than purely technical (Bibliometric Review of LLM Hallucination, 2025).

2.7 Socio-Technical Perspectives on Embedding GenAI into Strategy

A qualitative field study examining how a large multinational organization embedded generative AI into strategic decision-making processes, drawing on 27 semi-structured expert interviews conducted with decision-makers and strategy professionals, developed a process model identifying both enablers and barriers to successful integration through structured inductive coding (From Pilots to Decision Systems, 2025). Enablers identified included leadership-driven adoption, early quick-win projects, structured prompt-engineering experimentation, workforce training programs, secure deployment platforms, and dedicated investment allocation, while barriers included strategic ambiguity about generative AI's intended role, limited organizational awareness, hallucination risk, deficiencies in employee prompt-engineering capability, data readiness gaps, and unresolved privacy and intellectual property concerns (From Pilots to Decision Systems, 2025). This socio-technical and governance-lens study provides one of the more granular, process-level accounts of exactly the organizational learning gap the MIT NANDA report identifies at a larger, cross-sector scale.

3. Research Objectives and Question

This paper pursues four specific objectives: first, to synthesize current theoretical frameworks explaining generative AI's role in managerial and strategic decision-making; second, to assess the empirical reliability of AI-generated business judgments based on recent controlled studies; third, to map current enterprise application patterns and the structural barriers separating successful from unsuccessful deployments; and fourth, to synthesize emerging governance approaches addressing hallucination and reliability risk in high-stakes business decision contexts. The overarching research question is: what do current theoretical frameworks and empirical evidence reveal about generative AI's actual, as opposed to promised, contribution to business decision-making, and what structural and governance conditions determine whether that contribution translates into measurable enterprise value.

4. Research Methodology

This study uses a narrative literature review combined with secondary data synthesis, drawing on peer-reviewed journal articles, recent conference proceedings, and rigorously sourced industry research reports published primarily between 2023 and 2026. Given the extreme recency of generative AI as an organizational phenomenon, dating substantively from late 2022, the evidence base necessarily combines emerging academic literature, including several 2025 and 2026 journal articles and preprints, with large-sample industry survey research from established research organizations including MIT's NANDA initiative, McKinsey, Gartner, and Lucidworks; where a finding rests on industry research rather than peer-reviewed academic literature, this is indicated in the text. Sources were identified through targeted searches combining terms including generative AI, business decision-making, enterprise adoption, strategic management, hallucination, and AI governance, with priority given to studies employing empirical methods, controlled experiments, or large-sample surveys over purely conceptual or opinion-based commentary.

5. Findings and Discussion

5.1 Augmentation, Not Automation, Is the Empirically Supported Role

Across the theoretical and empirical literature reviewed, the strongest and most consistent finding is that generative AI currently functions reliably as a decision-augmentation tool rather than an autonomous decision-maker. The Strategic Management Journal evidence that aggregated, multi-model evaluations approximate human expert judgment while single-model evaluations do not (Doshi et al., 2025) provides a precise empirical boundary condition:

generative AI adds genuine value when its outputs are combined, cross-checked, and integrated with human judgment, and adds unreliable or even misleading value when a single AI output is trusted in isolation. This directly supports the augmentation-over-automation theoretical framing advanced by Toha et al. (2026) and the organizational decision-making review (Rethinking organizational decision-making, 2026), while simultaneously providing the kind of concrete, testable evidence that earlier, purely conceptual frameworks lacked.

5.2 Sycophancy and Ambiguity Blindness Are Underappreciated Failure Modes

The finding that generative AI systems frequently fail to recognize ambiguous business inputs and exhibit sycophantic agreement with flawed managerial directives (Generative AI in Managerial Decision-Making, 2026) deserves particular attention because it identifies a failure mode that is structurally different from, and arguably more dangerous than, simple factual hallucination. A hallucinated statistic can, in principle, be checked against a source. A sycophantic elaboration of a flawed strategic premise a manager has already committed to is much harder to detect, because it takes the surface form of competent, well-reasoned business analysis. This finding suggests that governance frameworks focused narrowly on factual accuracy and hallucination detection, however important, may be insufficient on their own, and that decision-support system design should explicitly test for and penalize excessive agreement with user-supplied premises.

5.3 The Adoption-Value Gap Is Organizational, Not Technical

The convergence between the MIT NANDA report's finding that the core barrier to scaling is a contextual "learning gap" rather than model capability, infrastructure, regulation, or talent (MIT NANDA, 2025), and the qualitative field study's granular finding that success depends on leadership-driven adoption, workforce training, and structured experimentation rather than model selection (From Pilots to Decision Systems, 2025), represents one of the more robust cross-methodological agreements in this literature. Two independent studies using entirely different methods, a large-sample industry survey and a 27-interview qualitative field study, arrive at the same structural diagnosis: enterprises are, in the large majority of cases, failing to realize generative AI's business decision-making value not because the underlying models are inadequate, but because organizational workflows, data readiness, and integration practices have not caught up with model capability.

5.4 The Partnership Advantage Has Direct Governance Implications

The finding that externally partnered generative AI implementations succeed at roughly triple the rate of purely internal builds, 67 percent versus 22 percent (MIT NANDA, 2025), has a direct bearing on how enterprises should structure their generative AI governance functions. Rather than treating internal build capability as inherently superior for reasons of security or customization, the evidence suggests enterprises should weight vendor partnership more heavily in build-versus-buy decisions specifically for generative AI, a conclusion that runs somewhat counter to conventional enterprise software governance instincts, which typically favor internal control for high-stakes systems.

5.5 Integration Across the Dimensions

Read together, the theoretical, empirical, applied, and governance literatures converge on an integrated picture: generative AI's demonstrated value in business decision-making is real but conditional, conditional on aggregating rather than trusting single outputs, on explicitly designing for ambiguity recognition and sycophancy resistance rather than assuming reliability, on prioritizing organizational learning capacity and workflow integration over model selection, and on governance architectures that calibrate risk response to decision consequentiality rather than pursuing unrealistic zero-hallucination targets. The 95-percent failure rate documented by MIT NANDA is best read not as evidence that generative AI lacks genuine decision-making value, since the Strategic Management Journal and socio-technical field-study evidence clearly demonstrate that value exists under the right conditions, but as evidence that most current enterprise implementations have not yet met those conditions.

6. Implications

For executives and strategy leaders: build explicit multi-model or multi-evaluation aggregation into any AI-assisted strategic decision process, rather than relying on single-model outputs, following the empirical pattern demonstrated by Doshi et al. (2025).

For AI governance functions: treat sycophancy and ambiguity-blindness testing as a required component of decision-support system validation, alongside conventional hallucination and bias testing, and calibrate governance rigor to decision consequentiality rather than applying uniform controls across low- and high-stakes use cases, following the risk-tiered approach proposed in the public-sector verification and validation literature (Risk-Oriented Verification and Validation Framework, 2026).

For technology and procurement leaders: weight external vendor partnership more heavily in generative AI build-versus-buy decisions given the demonstrated near-triple success rate advantage, while ensuring internal AI literacy and workforce training investment, since the socio-technical field-study evidence indicates vendor partnership alone is insufficient without corresponding internal organizational learning capacity (From Pilots to Decision Systems, 2025).

For researchers: the clearest gap in the current literature is the scarcity of studies that directly link specific governance and integration practices to measured business outcomes at scale; most current evidence either documents adoption patterns at scale, as in the MIT NANDA and McKinsey survey research, or documents decision-quality mechanisms in controlled or small-sample settings, as in the Strategic Management Journal and ambiguity-sycophancy studies, with comparatively little research connecting the two.

7. Limitations and Future Research

This synthesis is limited by the extreme recency of its subject matter: because generative AI's organizational deployment dates substantively from late 2022, most reviewed studies cover a relatively short observation window, and longitudinal evidence on whether the GenAI Divide narrows or widens over time is not yet available. Several sources reviewed here are industry research reports rather than peer-reviewed academic studies, and while the MIT NANDA report in particular has been widely cited and independently corroborated by other industry surveys, its specific methodology, 52 executive interviews and analysis of 300 public deployments, is narrower than a full academic sampling frame would typically require. Future research should prioritize longitudinal tracking of specific enterprise deployments to determine whether the structural barriers identified here, the learning gap, resource misallocation toward high-visibility functions, and the internal-build disadvantage, are transitional features of an early adoption phase or persistent structural features of generative AI integration. Future work should also extend the sycophancy and ambiguity-blindness findings from the single reviewed controlled study into larger-sample, cross-industry replications, given the significant governance implications these failure modes carry.

8. Conclusion

Generative artificial intelligence has achieved remarkably fast and broad adoption in business decision-making contexts, with a majority of organizations now using it in at least one business function within roughly three years of ChatGPT's public release. Yet the evidence reviewed in this paper shows that adoption breadth and decision-making value are currently poorly correlated: the same period that produced near-universal enterprise experimentation also produced robust evidence that 95 percent of pilots fail to deliver measurable financial impact. The theoretical and empirical literature converges on an explanation for this gap that is more organizational than technical, generative AI adds genuine, empirically demonstrated value to business decision-making when its outputs are aggregated rather than trusted singly, when systems are explicitly designed to resist ambiguity-blindness and sycophancy, and when organizations invest in the workflow integration and contextual learning capacity that separates the successful minority of deployments from the stalled majority. The central argument of this paper is that the current moment in generative AI and business decision-making should be understood not as a referendum on whether the technology works, but as an early, unevenly distributed test of whether organizations can build the governance, integration, and human-AI collaboration practices the technology's genuine capabilities actually require.

REFERENCES

1. Automely. (2026). The state of generative AI for business in 2026: Adoption, ROI, and what comes next.
2. Bastone, A., Mandiello, A., & Zeuli, F. (2026). Generative artificial intelligence and decision-making: Evidence from a participant observation with latent entrepreneurs. *European Journal of Innovation Management*. <https://doi.org/10.1108/EJIM-03-2025-0388>
3. Dong, S. C. (2025). Prudential reliability of large language models in reinsurance: Governance, assurance, and capital efficiency. *arXiv preprint*. <https://arxiv.org/pdf/2511.08082>
4. Doshi, A. R., Hauser, J., & Van den Bulte, C. (2025). Generative artificial intelligence and evaluating strategic decisions. *Strategic Management Journal*. <https://doi.org/10.1002/smj.3677>
5. eClerx. (2026). A practical guide for generative AI adoption for enterprises in 2026.
6. From Models to Metrics: A Governance Framework for Large Language Models in Enterprise AI and Analytics. (2026). [Special Issue: Critical Challenges in Large Language Models and Data Analytics]. <https://doi.org/10.3390/bdcc9010008>
7. From pilots to decision systems: Embedding generative AI into strategic decision-making through a socio-technical and governance lens. (2025). [Journal name pending verification]. <https://doi.org/10.1080/12460125.2025.2597835>
8. Generative AI in Managerial Decision-Making: Redefining Boundaries through Ambiguity Resolution and Sycophancy Analysis. (2026). *arXiv preprint*. <https://arxiv.org/html/2603.03970>
9. Generative Artificial Intelligence and Evaluating Strategic Decisions [conceptual study]. (2024). A Conceptual Study on the Effects of Artificial Intelligence in Managerial Decision-Making.
10. Generative AI Use in Entrepreneurship: An Integrative Review and an Empowerment-Entrapment Framework. (2026). *arXiv preprint*. <https://arxiv.org/pdf/2604.02567>
11. Kaplan, D. M., Palitsky, R., & Raison, C. L. (2025). The "machinal bypass" and how we're using AI to avoid ourselves. *Proceedings of the National Academy of Sciences*, 122(51), e2518999122. <https://doi.org/10.1073/pnas.2518999122>
12. Kim, J. (2025). Modeling generative AI and social entrepreneurial searches: A contextualized optimal stopping approach. *Administrative Sciences*, 15(8), 302. <https://doi.org/10.3390/admsci15080302>
13. Lucidworks. (2026). Enterprise AI in 2026: Adoption trends, gaps & strategic insights [2025 AI Benchmark Study].
14. MIT NANDA. (2025). The GenAI Divide: State of AI in business 2025.
15. Rethinking organizational decision-making: The emerging roles and tasks of generative artificial intelligence. (2026). *Management Review Quarterly*. <https://doi.org/10.1007/s11301-026-00611-2>
16. Risk-Oriented Verification and Validation Framework for Hallucination Detection in Public-Sector LLM Systems. (2026). [Journal name pending verification].
17. A Bibliometric Review of Large Language Model Hallucination. (2025). *International Journal of Research and Innovation in Social Science*.
18. LLM hallucination and bias detection in regulated enterprise systems. (2025). *World Journal of Advanced Research and Reviews*.
19. Toha, M. A., Khan, P. A., & Uddin, M. S. (2026). Theorizing generative AI's role in strategic decision-making: From automation to augmentation. In *AI, Society and Digital Transformation (ICSS 2025)*, Lecture Notes in Operations Research. Springer. https://doi.org/10.1007/978-3-032-13116-4_10
20. Simon, H. A. (1957). Models of man: Social and rational; mathematical essays on rational human behavior in society setting. Wiley.
21. Rai, A. (2020). Explainable AI: From black box to glass box. *Journal of the Academy of Marketing Science*, 48, 137–141.
22. Techment. (2025). Enterprise AI strategy in 2026: A proven roadmap for future-ready enterprises.
23. Larridin. (2026). AI adoption: The complete enterprise guide 2026 [Gartner survey data, 1,973 managers].
24. Stack AI. (2026). Enterprise AI adoption: State of generative AI in 2026