

AI-Augmented Management: Redefining Decision-Making, Leadership, and Organizational Performance

Qian Chen

European International University, Paris, France.

ABSTRACT

A factory manager in Ohio now receives a demand forecast from a machine learning model before her morning shift briefing. A hospital administrator in Singapore reviews an algorithmic staffing recommendation before approving next week's rota. Neither manager has been replaced. Both have gained a new kind of colleague, one that never tires of recalculating probabilities but cannot explain why a patient's readmission risk matters to the family in the waiting room. This paper examines how artificial intelligence is reshaping managerial work through augmentation rather than substitution, and asks what organizational conditions determine whether that augmentation strengthens or erodes decision quality, leadership effectiveness, and long-run performance. Drawing on the Resource-Based View, Dynamic Capabilities Theory, Upper Echelons Theory, Organizational Information Processing Theory, the Knowledge-Based View, Sociotechnical Systems Theory, Organizational Learning Theory, Human Capital Theory, and the Technology-Organization-Environment framework, the paper develops an integrated conceptual model in which AI capability, data quality, digital infrastructure, analytics capability, and AI governance act as antecedents to decision quality, knowledge sharing, organizational learning, leadership effectiveness, and employee empowerment, which in turn shape organizational performance, innovation capability, strategic agility, decision effectiveness, and competitive advantage. Organizational culture, digital maturity, environmental uncertainty, and ethical climate are proposed as boundary-setting moderators. The paper synthesizes existing empirical and conceptual literature, proposes nine testable propositions, outlines a mixed-methods survey design suited to structural equation modeling, and illustrates the framework across manufacturing, services, higher education, and healthcare settings. It closes with managerial, policy, and research implications, arguing that the durable competitive value of AI in management rests less on computational power than on the organizational, ethical, and human-capital conditions surrounding its use.

Keywords: Artificial intelligence, augmented decision-making, leadership effectiveness, organizational learning, AI governance, dynamic capabilities, strategic agility, organizational performance

INTRODUCTION

In March 2023, a mid-sized regional bank in the American Midwest began piloting a large language model to draft loan-committee memos. Within six months, the memos were faster to produce and more consistent in structure. What surprised the bank's chief credit officer was not the speed. It was that loan officers started asking sharper questions in committee meetings, because they no longer spent their preparation time formatting tables and instead spent it thinking about the borrower's business. That small, undocumented shift captures the central puzzle this paper takes up: artificial intelligence rarely replaces a manager's job in one clean motion. It reallocates attention, and what managers do with the attention they get back determines whether the technology pays off.

For much of the past decade, popular commentary on AI in the workplace has oscillated between two extremes: forecasts of mass managerial displacement and promises of frictionless organizational transformation. Neither extreme holds up well against the accumulating evidence. Raisch and Krakowski's influential analysis in the *Academy of Management Review* frames this as an automation-augmentation paradox: the same AI system can either narrow or expand human contribution, depending on how organizations choose to deploy it (Raisch & Krakowski, 2021). Davenport and Ronanki (2018), writing in the *Harvard Business Review*, reached a related conclusion from a practitioner survey of AI projects: the highest returns came not from replacing cognitive work but from augmenting it, particularly in tasks involving judgment under uncertainty.

This paper treats that observation as a starting point rather than a conclusion. If AI augments rather than replaces managerial cognition, then the interesting research questions shift away from job-loss projections and toward organizational design: which managerial responsibilities are strengthened, which decision processes change shape, what leadership competencies gain or lose relevance, and what governance arrangements keep the human judgment in human-AI collaboration meaningful rather than symbolic. These questions sit at the intersection of strategic management, organizational behavior, leadership studies, and management information systems, four fields that have historically developed AI-related insights somewhat independently of one another.

This paper sets out to build a bridge across that fragmentation. It develops an integrated conceptual framework that links AI-related organizational capabilities to decision quality, leadership effectiveness, and organizational learning, and traces their combined influence on performance outcomes including innovation capability, strategic agility, and competitive advantage. The framework draws on eight established theoretical traditions, each of which explains a different piece of the puzzle: why some firms build usable AI capability while others accumulate expensive dashboards nobody trusts, why some leaders extend their judgment with algorithmic input while others either defer to it uncritically or ignore it altogether, and why some organizations translate AI adoption into durable advantage while others see only temporary efficiency gains.

The remainder of the paper proceeds as follows. Section 2 reviews the literature on AI-augmented decision-making, leadership, and organizational capability, organized thematically rather than author by author. Section 3 contrasts optimistic and critical readings of that literature. Sections 4 through 9 identify the research gap, state objectives and questions, and lay out the theoretical foundations in full. Sections 10 through 13 present the integrated framework, nine propositions, and their boundary conditions. Section 14 details a methodology suited to testing the framework empirically. The remaining sections address expected contributions, managerial scenarios across four sectors, practical and policy implications, limitations, a future research agenda, and a closing synthesis.

LITERATURE REVIEW

From decision support to decision augmentation

Management scholars have studied computer-based decision aids since the 1970s, but the current wave of interest differs in an important respect: contemporary AI systems, particularly those built on machine learning and large language models, generate recommendations rather than merely organizing information for human review. A 2025 guest editorial introducing a special issue of *Management Decision* surveyed 24 accepted papers and concluded that AI systems tend to improve forecasting accuracy, speed up decision cycles, and raise the overall quality of managerial choices across a range of organizational contexts, from utilities to manufacturing (Massaro et al., 2025). That special issue also cautioned that these gains are unevenly distributed, appearing most reliably where firms had already invested in data infrastructure before adopting AI tools.

A parallel stream of work has examined generative AI specifically. A systematic review published in *Management Review Quarterly* identified 68 relevant studies and mapped 53 distinct generative AI tasks onto the classical stages of organizational decision-making, from attention and intelligence gathering through design, choice, implementation, and feedback (Schulte & Kanbach, 2026). The review's most useful contribution for this paper is its typology of AI roles: rather than treating AI as a single undifferentiated actor, it distinguishes between AI functioning as an information gatherer, an option generator, an evaluator, and, in a smaller but growing set of cases, a collaborative partner operating alongside a human decision-maker in what some scholars now call a metahuman system (Schulte & Kanbach, 2026).

Innovation management offers a useful test case because innovation decisions are famously resistant to routine, formulaic treatment. A systematic literature review published in *Technological Forecasting and Social Change* examined AI's role across the four canonical stages of the innovation process and found four distinct effects: AI can strengthen the rationality of innovation decisions, expand rather than constrain creative search, reorganize how innovation work gets coordinated, and introduce new categories of risk that earlier innovation-management theory did not anticipate (ScienceDirect, 2022). That last finding matters for this paper's later discussion of governance: augmentation is not a free good. It changes what can go wrong as much as it changes what can go right.

Leadership in the loop

If decision-making research asks what AI does to choices, leadership research asks what AI does to the people making them. A systematic literature review in the *Review of Managerial Science*, drawing on 63 articles from 31 journals and grounded explicitly in Upper Echelons Theory, found that senior leaders who actively reshape their own skills and mental models around AI tend to preside over more effective AI-related strategic decisions than leaders who delegate AI oversight downward without engaging with it themselves (Review of Managerial Science, 2025). This finding gives empirical weight to a long-standing claim in upper echelons research: strategic outcomes reflect the cognitive frames of top managers, and those frames are now being reshaped by exposure to algorithmic tools rather than only by industry experience or formal training.

A separate study using classical decision theory as its lens, including bounded rationality and Mintzberg's managerial role typology, examined how AI reshapes judgment, foresight, and leadership dynamics in international business settings and concluded that AI's clearest contribution is not to replace strategic foresight but to widen the range of scenarios a leader can consider before committing to a course of action (Kourkouvelis et al., 2024). Both studies converge on a claim this paper treats as foundational: leadership effectiveness in AI-augmented organizations depends less on technical fluency with the tools themselves and more on a leader's capacity to interrogate algorithmic output rather than accept or dismiss it wholesale.

Capability, knowledge, and organizational learning

A third stream connects AI adoption to organizational capability-building rather than to individual decisions or individual leaders. Work applying the Dynamic Capabilities View to AI adoption in human resource management, published in *Human Resource Management Review* (Journal of Business Research family), identified a specific set of managerial competencies needed to adopt AI-based tools responsibly, including the capacity to sense which processes are suitable for algorithmic support, to reconfigure workflows around new tools without discarding tacit expertise, and to build the internal trust needed for staff to actually use what has been deployed (ScienceDirect, 2021). This sensing-seizing-reconfiguring vocabulary maps directly onto Teece's original formulation of dynamic capabilities and gives this paper a bridge between classical strategy theory and contemporary AI-adoption research.

Knowledge management scholars have raised a related concern: AI systems trained on historical data can inadvertently encode and reproduce the organization's past blind spots, which threatens the double-loop learning that Argyris and Schön (1978) identified as the mechanism by which organizations correct their own operating assumptions rather than merely their outputs. If a firm uses AI chiefly to optimize existing routines, it may get faster at doing the wrong thing. This concern recurs throughout the governance literature discussed next and motivates this paper's inclusion of organizational learning as a mediating rather than purely dependent variable.

Governance, trust, and accountability

A fourth and increasingly active stream addresses the institutional conditions under which AI use in organizations remains accountable. Recent work distinguishes AI governance as a design problem, focused on transparency, auditability, and explainability, from AI governance as an institutional capacity problem, focused on whether an organization actually has the authority structures and technical visibility needed to exercise the oversight its governance documents promise (arXiv, 2026). This distinction matters for management research because much of the earlier literature assumed governance frameworks would function as intended once written down. A scoping review of trustworthiness scholarship across major AI ethics and fairness venues similarly found that organizational trust in AI systems is built incrementally through repeated, verifiable performance rather than through policy statements alone (arXiv, 2025). A review focused on ethical frameworks for organizational AI adoption reached a compatible conclusion: accountability requires that both AI developers and the organizational executives who deploy their systems share responsibility for outcomes, rather than treating accountability as something that can be outsourced entirely to a vendor's technical documentation (Madanchian & Taherdoost, 2025).

Human judgment, algorithmic recommendation, and the acceptance problem

A fifth theme cuts across the four already discussed: the question of when managers actually use the algorithmic recommendations available to them, as opposed to producing recommendations that sit unread in a dashboard. This is partly a question about system design and partly a question about organizational psychology. Executive cognition research suggests that senior managers weigh algorithmic input differently depending on whether a recommendation confirms or challenges their prior judgment, a pattern consistent with long-standing findings on confirmation bias in strategic decision-making rather than anything unique to AI. The practical consequence, raised repeatedly in the decision-augmentation literature reviewed above, is that AI adoption without attention to acceptance and trust tends to produce two failure modes at once: managers who defer to algorithmic output even when their own judgment would have served them better, and managers who ignore it even when it would have improved the decision. Both failure modes are more common in organizations that have not built the internal practices, structured second-opinion protocols, documented override rights, cross-functional review, that make human-AI collaboration a genuine partnership rather than a nominal one.

Cross-functional collaboration matters here in a way the leadership-focused literature in Section 2.2 does not fully capture on its own. AI-enabled strategic planning typically draws on data and interpretation from finance, operations, and frontline functions simultaneously, and the organizations that report the strongest results tend to be those that built cross-functional review into the planning cycle itself rather than treating AI output as a technical input handed to strategy teams after the fact. This observation connects the knowledge-sharing mechanism discussed in Section 2.3 to the leadership-engagement mechanism discussed in Section 2.2, and it is one of the reasons this paper treats knowledge sharing and leadership effectiveness as related but analytically separate mediators rather than collapsing them into a single construct.

Organizational resilience and the limits of efficiency gains

A smaller but growing body of work asks whether efficiency gains from AI adoption translate into organizational resilience, meaning the capacity to absorb shocks and adapt strategy under pressure, or whether the two are actually somewhat in tension. The concern raised most often is that AI systems optimized for stable historical patterns can perform poorly exactly when resilience matters most, during unprecedented disruptions where historical data offers limited guidance. This is not an argument against AI-augmented management so much as a caution against treating decision speed and forecasting accuracy, the two benefits most consistently documented in the literature reviewed in Section 2.1, as proxies for organizational resilience more broadly. The distinction supports this paper's decision to treat strategic agility and organizational performance as separate dependent variables rather than a single composite outcome.

Taken together, these six streams of literature point toward a common insight that existing studies have not yet assembled into a single testable model: AI augments managerial work through a chain of organizational conditions, not through the technology alone. Capability, leadership engagement, learning processes, and governance arrangements each act as a gate through which the benefits of AI must pass before they show up in organizational performance.

A thematic map of the literature

Read side by side, the six streams reviewed above cluster around three questions rather than six independent topics. The first question, addressed in Sections 2.1 and 2.5, is whether a given decision improves when AI enters the process, and the answer depends heavily on data quality, task uncertainty, and whether managers treat algorithmic output as one input among several rather than a verdict. The second question, addressed in Sections 2.2 and 2.6, is whether the people responsible for decisions change how they lead as a result of working alongside AI, and the answer depends on whether senior leaders engage with the tools directly or delegate that engagement away. The third question, addressed in Sections 2.3 and 2.4, is whether the organization as a whole becomes more capable and more accountable over time, and the answer depends on capability-building and governance capacity rather than on any single AI deployment. These three questions correspond closely to the decision quality, leadership effectiveness, and combined learning and governance mechanisms that anchor the conceptual framework developed in Sections 8 and 9, which is by design: the framework is an attempt to give this fragmented map a shared structure rather than an independent theoretical invention layered on top of it.

Sections 6 through 11 of this paper build that map into an explicit conceptual framework.

COMPETING PERSPECTIVES IN THE LITERATURE

Two broad readings of this literature deserve to be stated side by side rather than blended together.

The optimistic reading, most closely associated with the decision-support and dynamic-capabilities streams reviewed above, holds that AI expands the effective cognitive capacity of management teams. On this view, tasks that once consumed scarce managerial attention, pattern recognition across large datasets, routine forecasting, and first-pass option generation, can be handed to algorithmic systems, freeing human judgment for exactly the tasks it is best suited for: weighing values, handling ambiguity, and taking responsibility for outcomes. Proponents of this view point to documented gains in forecasting accuracy and decision speed (Massaro et al., 2025) as early evidence that the augmentation thesis is more than aspirational rhetoric.

The critical reading, most visible in the governance and accountability literature, treats these same gains with greater caution. Writers in this tradition point out that decision speed is not the same as decision

quality, that algorithmic recommendations can appear more objective than they are because their statistical basis is opaque to the managers using them, and that organizations frequently adopt AI governance language faster than they build the institutional capacity to back it up (arXiv, 2026). A related concern, raised in the innovation-management literature, is that AI can narrow rather than widen the space of options a decision-maker considers, particularly when algorithmic recommendations are treated as a default rather than as one input among several (ScienceDirect, 2022).

Both readings are compatible with the empirical record; they simply emphasize different boundary conditions. The optimistic account tends to describe organizations with mature data infrastructure, engaged leadership, and functioning governance. The critical account tends to describe organizations that adopted AI tools without those conditions in place. This paper's conceptual framework treats that divergence as the central empirical question rather than resolving it by assumption, which is why organizational culture, digital maturity, environmental uncertainty, and ethical climate appear as moderators rather than as background variables.

RESEARCH GAP

Three gaps in the existing literature motivate this paper. First, the decision-making, leadership, capability, and governance literatures reviewed above have developed largely in parallel, each publishing in its own set of journals and citing a largely separate set of foundational theories. No existing model links AI capability, data quality, digital infrastructure, analytics capability, and AI governance to organizational performance through the combined mediating mechanisms of decision quality, knowledge sharing, organizational learning, leadership effectiveness, and employee empowerment. Second, much of the empirical literature to date has measured AI adoption as a binary or intensity variable without accounting for the organizational conditions, such as ethical climate and digital maturity, under which the same adoption produces different outcomes. Third, sector-specific applications remain scattered across disconnected case studies rather than organized within a shared conceptual structure that would allow comparison across manufacturing, services, education, and healthcare settings. This paper addresses each of these gaps through the integrated framework developed in Sections 10 through 12.

RESEARCH OBJECTIVES

This paper pursues five objectives: to trace the evolution of AI's role in management from decision support toward decision augmentation; to analyze how AI-supported decision processes differ from traditional managerial decision processes; to examine how AI adoption is reshaping the competencies associated with effective leadership; to develop and justify an integrated conceptual framework linking AI-related organizational capability to performance outcomes through identified mediating and moderating mechanisms; and to translate that framework into recommendations for managers and policymakers responsible for governing AI use inside organizations.

RESEARCH QUESTIONS

The paper is organized around six questions: How is AI transforming managerial decision-making? Which managerial responsibilities are enhanced rather than automated by AI? How does AI influence leadership effectiveness and organizational learning? What organizational capabilities determine successful AI-augmented management? What governance mechanisms are necessary to ensure responsible managerial use of AI? How does AI contribute to sustainable organizational performance?

THEORETICAL FOUNDATIONS

Building an integrated framework from eight distinct theoretical traditions requires being explicit about what each tradition contributes and where its explanatory reach ends.

Resource-Based View. Barney's (1991) argument that sustained competitive advantage arises from resources that are valuable, rare, imperfectly imitable, and non-substitutable supplies this paper's starting premise: AI infrastructure purchased off the shelf is rarely a source of advantage on its own, because competitors can buy the same infrastructure. What resists imitation is the organizationally embedded capability to use that infrastructure well, which is why this paper treats AI capability, rather than AI technology as such, as the relevant independent variable.

Dynamic Capabilities Theory. Teece, Pisano, and Shuen's (1997) framework of sensing, seizing, and reconfiguring extends the Resource-Based View into environments where conditions change quickly, which describes the AI landscape reasonably well. Organizations need the capacity to sense which AI applications fit their decision problems, to seize opportunities by integrating tools into workflows, and to reconfigure routines as tools and data sources evolve. This theory grounds the paper's treatment of AI capability and analytics capability as dynamic rather than fixed resources.

Upper Echelons Theory. Hambrick and Mason's (1984) claim that organizational outcomes reflect the characteristics, cognitions, and values of top managers explains why leadership engagement, rather than technology procurement alone, predicts AI-related performance gains in the empirical literature reviewed in Section 2.2. This theory anchors leadership effectiveness as a mediating variable in the conceptual model.

Organizational Information Processing Theory. Galbraith's (1974) account of organizations as information-processing systems that must match processing capacity to task uncertainty offers a natural lens for AI adoption, since AI systems primarily change an organization's capacity to process information under uncertainty. This theory supports the inclusion of environmental uncertainty as a moderator: the value of added information-processing capacity should be greatest precisely where uncertainty is highest.

Knowledge-Based View. Grant's (1996) extension of the Resource-Based View toward knowledge as the most strategically significant organizational resource informs this paper's treatment of knowledge sharing as a mediating mechanism. AI systems can either widen access to organizational knowledge, by surfacing patterns that would otherwise remain tacit, or narrow it, by concentrating interpretive authority in whoever controls the algorithmic tools.

Sociotechnical Systems Theory. Trist and Bamforth's (1951) foundational insight, that organizational performance depends on jointly optimizing technical and social systems rather than optimizing either alone, is arguably the single most directly relevant theory in this list. It cautions against treating AI deployment as a purely technical project and supports this paper's inclusion of organizational culture and ethical climate as moderating variables that condition whether technical capability translates into performance.

Organizational Learning Theory. Argyris and Schön's (1978) distinction between single-loop learning, which corrects errors within existing assumptions, and double-loop learning, which questions the assumptions themselves, is essential for understanding the risk raised in Section 2.3: AI systems trained on historical data can reinforce single-loop correction while suppressing the double-loop questioning that produces genuine strategic adaptation. Organizational learning therefore appears in the model as a mediator whose direction of effect is not guaranteed to be positive.

Human Capital Theory. Becker's (1964) framework linking investment in skills and knowledge to individual and organizational productivity supports this paper's treatment of employee empowerment as a mediating variable distinct from leadership effectiveness. AI adoption changes the return on

different kinds of human capital, generally raising the value of judgment, interpretation, and interpersonal skill relative to the value of routine information-processing skill.

Technology-Organization-Environment Framework. Tornatzky and Fleischer's (1990) model, which explains technology adoption as a function of technological characteristics, organizational characteristics, and environmental context, is used here less as a standalone theory and more as an organizing device that groups the model's four categories of variables, independent, mediating, moderating, and dependent, into a coherent adoption-to-performance narrative.

No single theory listed above is sufficient on its own to explain AI-augmented management, and this is itself an argument for the integrative approach this paper takes. The Resource-Based View and Dynamic Capabilities Theory explain why some firms build usable AI capability, but neither says much about what happens inside a leader's head once that capability exists, which is where Upper Echelons Theory contributes. Organizational Information Processing Theory explains why the value of AI-driven information capacity varies with task uncertainty, but it does not address the social conditions under which employees actually trust and use that capacity, which is where Sociotechnical Systems Theory and the Knowledge-Based View contribute. Organizational Learning Theory and Human Capital Theory then explain what happens over time, whether the organization's stock of knowledge and skill deepens or stagnates as AI use becomes routine. The Technology-Organization-Environment framework, finally, supplies the structural scaffolding that lets these otherwise separate contributions sit inside one coherent model rather than eight competing ones.

INTEGRATED CONCEPTUAL FRAMEWORK

The framework proposed here treats AI Capability, Data Quality, Digital Infrastructure, Analytics Capability, and AI Governance as independent variables. These five factors are not treated as interchangeable; AI capability captures the organization's dynamic capacity to sense and deploy AI tools, data quality captures the reliability and representativeness of the information those tools consume, digital infrastructure captures the technical backbone that makes deployment possible, analytics capability captures the human and procedural skill needed to interpret model output, and AI governance captures the accountability structures that keep deployment aligned with organizational values.

These five independent variables act on organizational performance through five mediating mechanisms: Decision Quality, Knowledge Sharing, Organizational Learning, Leadership Effectiveness, and Employee Empowerment. None of the five independent variables is expected to affect performance directly without passing through at least one of these mediators; a firm can own excellent AI infrastructure and see no performance benefit if that infrastructure never translates into better decisions, wider knowledge sharing, more effective leadership, more genuine organizational learning, or a more empowered workforce.

This full-mediation assumption is deliberate rather than a simplifying convenience. Much of the earlier technology-adoption literature modeled performance as a fairly direct function of adoption intensity, and the AI-specific findings reviewed in Section 2 suggest that assumption no longer fits well. Firms adopt broadly similar AI tools and report widely divergent results, which is difficult to explain unless the mediating mechanisms, not the technology itself, are doing most of the causal work. Treating the five mediators as the primary carriers of AI's organizational effect also gives the framework practical value beyond its research use: it gives managers five specific, actionable places to look when an AI investment underperforms, rather than a single diffuse question about whether the technology was worth buying.

Four moderators condition the strength of these mediated relationships: Organizational Culture, Digital Maturity, Environmental Uncertainty, and Ethical Climate. Consistent with sociotechnical systems

theory, these moderators are treated as necessary rather than incidental features of the model. Finally, the framework proposes five dependent variables: Organizational Performance, Innovation Capability, Strategic Agility, Decision Effectiveness, and Competitive Advantage. These are kept analytically distinct because the literature reviewed in Section 2 suggests AI's effects on them are not uniform; a firm may see rapid gains in decision effectiveness well before those gains consolidate into durable competitive advantage.

PROPOSITION DEVELOPMENT

Rather than stating isolated hypotheses, the propositions below are developed as a connected sequence, each building on theoretical grounds established in Section 7.

Proposition 1. Organizations with higher AI capability, understood as the dynamic capacity to sense, integrate, and reconfigure AI tools, will exhibit higher decision quality than organizations that adopt equivalent AI technology without corresponding capability, consistent with the Resource-Based View's distinction between owning a resource and being able to use it well.

Proposition 2. Data quality will moderate the relationship between AI capability and decision quality, such that the benefits of AI capability are attenuated in data environments characterized by incompleteness, bias, or poor representativeness, in line with Organizational Information Processing Theory's requirement that processing capacity be matched to the reliability of the information being processed.

Proposition 3. Analytics capability will mediate the relationship between digital infrastructure and knowledge sharing, because infrastructure alone does not determine whether insights generated by AI systems are interpreted correctly and distributed across the organization rather than confined to technical specialists.

Proposition 4. AI governance will be positively associated with organizational trust in AI-supported decisions, and this relationship will be stronger where governance reflects genuine institutional capacity for oversight rather than documentation alone, consistent with recent governance scholarship distinguishing design-level from capacity-level accountability.

Proposition 5. Leadership effectiveness will mediate the relationship between AI capability and strategic agility, such that AI capability enhances strategic agility primarily through leaders who actively reshape their own decision frames around algorithmic input, rather than through the technology's direct effect on organizational routines, consistent with Upper Echelons Theory.

Proposition 6. Organizational learning will mediate the relationship between knowledge sharing and innovation capability, but this mediation will be positive only where AI-supported knowledge sharing supports double-loop rather than single-loop learning, in line with Argyris and Schön's distinction.

Proposition 7. Employee empowerment will mediate the relationship between AI governance and decision effectiveness, because governance arrangements that clarify accountability tend to increase employees' willingness to exercise independent judgment alongside AI recommendations rather than defer to them uncritically.

Proposition 8. Ethical climate will moderate the relationship between AI governance and employee empowerment, such that formal governance structures produce genuine empowerment only in organizations where ethical climate supports raising concerns about AI-supported decisions without professional risk.

Proposition 9. The combined mediating chain, decision quality, knowledge sharing, organizational learning, leadership effectiveness, and employee empowerment, will jointly predict competitive advantage more strongly than any single mediator considered alone, reflecting this paper's central claim that AI's organizational value is distributed across multiple capability domains rather than concentrated in any one of them.

BOUNDARY CONDITIONS

The framework developed above is not proposed as a universal account of AI's organizational effects, and several boundary conditions deserve explicit statement. First, the model assumes an organizational context with some baseline level of digital infrastructure; in settings where basic digital record-keeping is absent, the independent variables as specified may not apply in any meaningful sense, and adoption questions become infrastructural rather than managerial. Second, the model is built primarily from evidence in knowledge-intensive organizations, and its applicability to labor-intensive or highly routinized production environments, where AI's role may be closer to automation than augmentation, is likely to be weaker. Third, the model assumes AI systems are deployed with some degree of managerial discretion over use; in heavily regulated or algorithm-mandated environments, such as certain credit-underwriting or clinical-triage contexts, the moderating role of organizational culture may be constrained by external compliance requirements that leave less room for local adaptation. Finally, the propositions assume a relatively stable AI system over the observation period; because underlying models and their training data continue to change, longitudinal tests of this framework will need to account for the technology itself evolving during the study window, a complication not present in most earlier organizational technology research.

RESEARCH METHODOLOGY

Research philosophy. A post-positivist philosophy is appropriate for this study, given its aim of testing theoretically derived relationships between measurable constructs while acknowledging that organizational phenomena, including managerial judgment and leadership effectiveness, cannot be captured with the same precision as physical measurements.

Research approach. A deductive approach is proposed, moving from the theoretically grounded propositions in Section 9 toward hypothesis testing, complemented by a qualitative component intended to contextualize quantitative findings rather than to generate new theory from scratch.

Research design. A concurrent mixed-methods design is recommended: a cross-sectional quantitative survey administered to managers in AI-adopting organizations, paired with semi-structured interviews conducted with a smaller subsample to capture the interpretive detail that survey instruments tend to miss, particularly around governance practice and leadership sensemaking.

Target population. The target population consists of managers and executives at the middle and senior levels in knowledge-intensive organizations, defined as organizations in which a substantial share of value creation depends on information processing and judgment rather than on physical production alone. This includes financial services, professional services, higher education administration, and healthcare management, alongside knowledge-intensive functions within manufacturing firms.

Sampling strategy. A stratified purposive sampling strategy is recommended, stratifying by sector (manufacturing, services, higher education, healthcare) and by organizational AI maturity (informed by prior self-reported adoption stage), to ensure the sample captures variation on key moderators rather than concentrating in early-adopter organizations, which tend to be overrepresented in convenience samples drawn from technology-conference attendee lists.

Sample size justification. For a structural model of the complexity proposed here, with five independent variables, five mediators, four moderators, and five dependent variables, a minimum

sample of 300 to 400 respondents is recommended, consistent with common rules of thumb for covariance-based structural equation modeling requiring roughly ten to twenty observations per estimated parameter, and consistent with sample sizes typical of comparable published studies in this literature (Khan et al., 2026).

Measurement scales. Existing validated scales should be adapted rather than created from scratch wherever possible: AI capability and analytics capability can draw on established dynamic-capability measurement scales adapted to the AI context; decision quality can draw on established multi-item decision-effectiveness scales; leadership effectiveness can draw on established transformational and adaptive leadership scales; organizational learning can draw on established single-loop and double-loop learning scales derived from Argyris and Schön's framework; and AI governance and ethical climate can draw on recently published AI-specific governance and trust scales emerging from the accountability literature reviewed in Section 2.4. All items should use seven-point Likert response formats for consistency with standard structural equation modeling practice.

Questionnaire constructs. The instrument should include a screening section confirming respondent role and organizational AI adoption status, followed by separate blocks for each of the fourteen constructs in the model (five independent, five mediating, four moderating), a demographic and firmographic block, and a final open-response item inviting brief qualitative elaboration to support the interview phase.

Reliability assessment. Internal consistency should be assessed using Cronbach's alpha and composite reliability, with a threshold of 0.70 or above for each construct, alongside average variance extracted for convergent validity.

Construct validity. Confirmatory factor analysis should precede structural testing to establish that indicators load appropriately on their intended latent constructs, with discriminant validity assessed using the Fornell-Larcker criterion and the heterotrait-monotrait ratio.

Common method bias assessment. Because the survey relies on self-reported managerial perceptions collected at a single point in time, Harman's single-factor test and a marker-variable technique are recommended, alongside procedural remedies such as separating predictor and outcome items within the questionnaire and assuring respondent anonymity.

Ethical considerations. The study requires institutional ethics review, informed consent, anonymized data handling, and voluntary participation with the right to withdraw, along with particular care around questions that ask respondents to characterize their own organization's governance failures, which should be framed to minimize professional risk to participants.

Statistical techniques. Partial Least Squares Structural Equation Modeling (PLS-SEM) is recommended over covariance-based SEM given the model's exploratory elements and its combination of formative and reflective constructs, though Confirmatory Factor Analysis should still be used to validate the measurement model before structural testing. Mediation should be tested using bootstrapped indirect effects, and moderation should be tested using interaction terms with mean-centered variables. Robustness checks should include multi-group analysis across the four sector strata to test whether the model's structural paths hold consistently across manufacturing, services, higher education, and healthcare, and a comparison of alternative model specifications to rule out equally plausible causal orderings among the mediators.

EXPECTED CONTRIBUTIONS

This paper's theoretical contribution lies in integrating eight previously separate theoretical traditions into a single testable model of AI-augmented management, something the fragmented literature reviewed in Section 2 has not yet accomplished. Its managerial contribution lies in giving practicing managers a structured way to diagnose why AI investment sometimes pays off and sometimes does not, by directing attention to capability, governance, and culture rather than to technology procurement alone. Its methodological contribution lies in proposing a measurement approach, built from adapted existing scales, that could be applied consistently across sectors to make future studies comparable in ways the current scattered case-study literature does not allow. Its policy relevance lies in distinguishing governance-as-documentation from governance-as-institutional-capacity, a distinction regulators and standard-setting bodies are only beginning to build into formal requirements. Its contribution to emerging management research lies in treating organizational learning and leadership effectiveness as mediators with uncertain sign, rather than assuming AI adoption automatically improves them, which opens space for future studies to examine conditions under which AI use degrades rather than strengthens organizational learning.

MANAGERIAL SCENARIOS

Manufacturing. A plant manager overseeing predictive-maintenance AI faces a decision quality question with unusually clear stakes: false negatives risk equipment failure, false positives waste maintenance budget. The framework predicts that analytics capability, the shop floor's capacity to interpret model confidence intervals rather than treat alerts as binary commands, will determine whether predictive maintenance improves decision effectiveness or simply shifts the burden of judgment onto technicians without giving them the tools to exercise it. A secondary implication concerns knowledge sharing across shifts: plants that build routines for technicians to feed their own field observations back into the model, rather than treating the system as a one-way source of instructions, are more likely to see the double-loop learning benefits described in Proposition 6, because those observations often reveal maintenance patterns the training data never captured.

Services. A regional bank's loan committee, similar to the example opening this paper, illustrates leadership effectiveness as a mediator. The framework predicts that committee chairs who model active interrogation of AI-generated credit memos, rather than either rubber-stamping or dismissing them, will produce measurably better lending decisions than chairs who adopt either extreme posture, consistent with Proposition 5. The scenario also illustrates cross-functional collaboration in practice: credit memos increasingly draw on data assembled jointly by risk, relationship management, and compliance staff, and committees that formalize this joint assembly into the review cycle tend to catch data-quality problems, an outdated income figure, a stale collateral valuation, before those problems reach the point of decision, which is consistent with Proposition 2's prediction about data quality moderating decision quality.

Higher education. A university admissions office using AI to screen applications faces an ethical climate test directly. The framework predicts that admissions staff will be more willing to override algorithmic recommendations, and will do so more accurately, in departments where raising concerns about a recommendation carries no professional cost, consistent with Proposition 8's link between ethical climate and employee empowerment.

Healthcare. A hospital deploying AI-based readmission-risk scoring illustrates the boundary condition discussed in Section 10. Clinical decisions are heavily shaped by regulatory and professional standards that constrain local managerial discretion, so the framework predicts that AI governance in this setting will matter more for legal and reputational risk management than for the innovation-capability pathway that dominates in less regulated sectors. At the same time, the administrator's own experience

described at the start of this paper suggests that decision quality still depends on interpretive work that governance policy alone cannot supply: a risk score attached to a discharge plan means little until a case manager translates it into a conversation about home support, transportation, and the family's own capacity to manage a recovery, work that sits squarely in the domain of employee empowerment rather than algorithmic output.

PRACTICAL IMPLICATIONS

Organizations seeking to capture the performance benefits documented in the literature reviewed above should treat AI adoption as an organizational design project rather than a procurement decision. Investment in AI capability and analytics capability should precede or accompany investment in AI technology itself, since the empirical record consistently shows technology without capability producing weaker outcomes. Leadership development programs should explicitly address how senior managers engage with algorithmic recommendations, since Upper Echelons Theory and the supporting empirical evidence both suggest that leadership engagement, not technical literacy alone, predicts strategic benefit. Finally, organizations should treat governance as an ongoing institutional practice rather than a one-time policy document, building the internal capacity to actually exercise the oversight their governance frameworks promise.

A second practical implication concerns employee acceptance, a topic that recurs throughout the literature reviewed in Section 2 but rarely receives dedicated managerial attention during implementation planning. Acceptance is not simply a matter of training employees to operate new software. The evidence summarized in Sections 2.5 and 2.6 suggests acceptance depends on whether employees believe they retain genuine authority to override algorithmic recommendations and whether the organization's culture treats such overrides as a normal part of professional judgment rather than as a sign of distrust in company technology. Managers rolling out AI tools would do well to build override tracking into the system from the outset, not to penalize disagreement with the model but to create a record that lets the organization learn from cases where human judgment and algorithmic recommendation diverged, since those cases are often the most informative ones available.

A third implication concerns middle management specifically, a group largely absent from the leadership literature reviewed in Section 2.2, which tends to focus on senior executives. Middle managers frequently sit at the point where AI-generated recommendations meet frontline implementation, and the framework developed here suggests their capacity to translate algorithmic output into locally meaningful guidance functions as a distinct empowerment mechanism worth attention in its own right. Organizations that invest leadership development resources exclusively at the executive level risk leaving this translation function under-resourced, even where senior leadership engagement with AI is otherwise strong.

POLICY IMPLICATIONS

For policymakers and standard-setting bodies, this paper's distinction between governance-as-documentation and governance-as-capacity carries a direct implication: regulatory requirements that mandate written AI governance policies without verifying organizational capacity to implement them risk producing compliance activity that does little to protect the outcomes those policies target. Sector regulators, particularly in finance, healthcare, and education, should consider requiring evidence of institutional capacity, staffing, escalation pathways, audit access, alongside policy documentation. Public funding bodies supporting AI adoption in public-sector organizations, including universities and hospitals, should similarly attach capability-building requirements to technology grants rather than funding infrastructure procurement alone.

There is also a labor-market dimension to policy that the management literature reviewed here touches only indirectly. If the augmentation thesis central to this paper is correct, and AI raises the return to judgment, interpretation, and interpersonal skill relative to routine information processing, then workforce development policy built around technical AI literacy alone addresses only part of the transition workers face. Human Capital Theory, discussed in Section 7, would predict that the more durable policy investment is in the judgment-intensive skills that complement rather than compete with algorithmic tools, a point with implications for vocational training curricula, professional accreditation standards, and public employment services well beyond the organizations examined directly in this paper.

LIMITATIONS

This paper is conceptual rather than empirical, and its propositions await the testing outlined in Section 11 before they can be treated as established findings. The theoretical integration attempted here necessarily simplifies each of the eight source theories to fit a shared framework, and specialists in any one tradition may reasonably find that simplification limiting. The literature reviewed draws heavily on recent publications, some quite recent, and the underlying technology continues to change at a pace that may outstrip the rate at which supporting empirical evidence accumulates. Finally, the managerial scenarios in Section 13 are illustrative rather than based on original fieldwork, and should be read as hypothesis-generating rather than as evidence in their own right.

FUTURE RESEARCH AGENDA

Four directions for future research follow directly from the gaps identified above. First, longitudinal studies are needed to test whether the mediating relationships proposed here hold as underlying AI systems themselves change over the study period, a methodological challenge distinct from most earlier organizational technology research. Second, comparative studies across the four sectors discussed in Section 13 would help establish whether the framework's structural paths genuinely generalize or whether sector-specific models are required. Third, research directly measuring whether AI-supported knowledge sharing produces single-loop or double-loop learning, rather than assuming the more favorable case, would address a genuine uncertainty in the current literature. Fourth, research on the lived experience of managers navigating AI governance failures, rather than governance frameworks considered abstractly, would strengthen the policy implications proposed in Section 15.

CONCLUSION

The central claim of this paper is a modest one stated plainly: artificial intelligence changes what managers pay attention to, and organizations that treat this as a capability-building problem tend to do better than organizations that treat it as a purchasing decision. The eight theoretical traditions integrated here, spanning strategy, organizational behavior, leadership, and information systems, converge on that claim from different directions, which is itself evidence that the claim reflects something real about how organizations absorb new information-processing technology rather than an artifact of any single theoretical lens. The nine propositions developed in Section 9 translate that claim into testable form, and the methodology outlined in Section 11 offers a concrete path toward doing so. Whether AI ultimately strengthens or narrows managerial judgment will depend less on the sophistication of the models organizations deploy than on the governance, culture, and leadership engagement surrounding their use, the same organizational conditions that have determined the fate of nearly every major management technology before it.

REFERENCES

Argyris, C., & Schön, D. A. (1978). *Organizational learning: A theory of action perspective*. Addison-Wesley.

- Barney, J. (1991). Firm resources and sustained competitive advantage. *Journal of Management*, 17(1), 99-120. <https://doi.org/10.1177/014920639101700108>
- Becker, G. S. (1964). *Human capital: A theoretical and empirical analysis, with special reference to education*. University of Chicago Press.
- Davenport, T. H., & Ronanki, R. (2018). Artificial intelligence for the real world. *Harvard Business Review*, 96(1), 108-116.
- Massaro M, Secinaro S, Bagnoli C, Calandra D (2025), "Guest editorial: Artificial Intelligence (AI) for management decision-making processes. From measurement to strategy". *Management Decision*, Vol. 63 No. 10 pp. 3213–3225, <https://doi.org/10.1108/MD-10-2025-377>
- Galbraith, J. R. (1974). Organization design: An information processing view. *Interfaces*, 4(3), 28-36. <https://doi.org/10.1287/inte.4.3.28>
- Grant, R. M. (1996). Toward a knowledge-based theory of the firm. *Strategic Management Journal*, 17(S2), 109-122. <https://doi.org/10.1002/smj.4250171110>
- Hambrick, D. C., & Mason, P. A. (1984). Upper echelons: The organization as a reflection of its top managers. *Academy of Management Review*, 9(2), 193-206. <https://doi.org/10.5465/amr.1984.4277628>
- Khan, I., Gul, M., & Shah, S. M. (2026). The impact of artificial intelligence adoption on managerial decision-making effectiveness. *Journal of Management Science Research Review*, 5(1), 1196–1218.
- Schulte, N., & Kanbach, D. K. (2026). Rethinking organizational decision-making: The emerging roles and tasks of generative artificial intelligence. *Management Review Quarterly*. Advance online publication. <https://doi.org/10.1007/s11301-026-00611-2>
- Madanchian, M., & Taherdoost, H. (2025). Ethical theories, governance models, and strategic frameworks for responsible AI adoption and organizational success. *Frontiers in Artificial Intelligence*, 8, 1619029. <https://doi.org/10.3389/frai.2025.1619029>
- Raisch, S., & Krakowski, S. (2021). Artificial intelligence and management: The automation-augmentation paradox. *Academy of Management Review*, 46(1), 192-210. <https://doi.org/10.5465/amr.2018.0072>
- Kourkouvelis, A., Anastasopoulou, E. E., Deirmentzoglou, G. A., Masouras, A., & Anastasiadou, S. (2024). Artificial intelligence and managerial decision-making in international business. *European Conference on Innovation and Entrepreneurship*, 19(1), 363–369. <https://doi.org/10.34190/ecie.19.1.2573>
- Review of Managerial Science. (2025). Enhancing top managers' leadership with artificial intelligence: Insights from a systematic literature review. *Review of Managerial Science*. <https://doi.org/10.1007/s11846-025-00836-7>
- ScienceDirect. (2021). Decision augmentation and automation with artificial intelligence: Threat or opportunity for managers? <https://www.sciencedirect.com/science/article/abs/pii/S0007681321000288>
- ScienceDirect. (2022). A solution looking for problems? A systematic literature review of the rationalizing influence of artificial intelligence on decision-making in innovation management. *Technological Forecasting and Social Change*. <https://www.sciencedirect.com/science/article/abs/pii/S0040162522003523>
- Teece, D. J., Pisano, G., & Shuen, A. (1997). Dynamic capabilities and strategic management. *Strategic Management Journal*, 18(7), 509-533. [https://doi.org/10.1002/\(SICI\)1097-0266\(199708\)18:7%3C509::AID-SMJ882%3E3.0.CO;2-Z](https://doi.org/10.1002/(SICI)1097-0266(199708)18:7%3C509::AID-SMJ882%3E3.0.CO;2-Z)
- Tornatzky, L. G., & Fleischer, M. (1990). *The processes of technological innovation*. Lexington Books.
- Trist, E. L., & Bamforth, K. W. (1951). Some social and psychological consequences of the longwall method of coal-getting. *Human Relations*, 4(1), 3-38. <https://doi.org/10.1177/001872675100400101>