



Paper Type: Original Article

From Predictive Analytics to Agentic Decision Intelligence: A Governance-Aware Framework for Human–AI Management Systems

Abbas Shekari¹, Rajabali Ghasem Pour^{2,*} 

¹ Department of Industrial Engineering Department, Yazd university, Yazd, Iran; a_shekari@yahoo.com.

² Department of Mathematics and Physics, University of Campania “Luigi Vanvitelli”, Caserta, Italy; rajabali.ghasempour@studenti.unicampania.it.

Citation:

Received: 22 July 2026

Revised: 13 September 2026

Accepted: 21 November 2026

Shekari, A., & Ghasem Pour, R. (2026). From predictive analytics to agentic decision intelligence: A Governance-aware framework for Human–AI management systems. *Management Analytics and Social Insights*, 3(3), 244–257.

Abstract

Organizations are moving beyond dashboards and predictive models toward Artificial Intelligence (AI) systems that can interpret goals, plan multi-step work, invoke digital tools, and initiate actions. Yet management research still lacks a coherent model for deciding when such autonomy creates value and when it creates unacceptable organizational risk. This conceptual article develops the construct of Agentic Decision Intelligence (ADI) to describe governance-bounded AI-enabled decision systems that combine analytical inference, Large Language Model (LLM) reasoning, organizational knowledge retrieval, tool use, and monitored action. Drawing on research in business analytics, human–AI collaboration, algorithmic management, explainability, trust, autonomous agents, and AI governance, the paper develops a five-layer architecture linking organizational sensing, evidence construction, decision reasoning, controlled execution, and learning. It then introduces a decision-rights model that assigns human and machine authority according to consequence severity, environmental uncertainty, reversibility, and evidentiary strength. Four governance mechanisms, provenance, permissioning, escalation, and continuous assurance, are integrated into the architecture so that autonomy becomes conditional rather than binary. The resulting framework explains why agentic systems should not be evaluated only by predictive accuracy or task completion: managerial value depends on the joint quality of decisions, actions, accountability, and organizational learning. The article contributes a vocabulary for distinguishing analytical, generative, and agentic capabilities; a governance-aware architecture for management analytics; six testable propositions; and a staged implementation roadmap. These contributions reposition management analytics from insight production toward accountable decision execution and provide a research agenda for designing human–AI systems that are both operationally useful and institutionally governable.

Keywords: Agentic artificial intelligence, Management analytics, Decision intelligence, Human–AI collaboration, AI governance, Organizational decision-making.

1 | Introduction

Management analytics has historically sought to improve organizational decisions by transforming data into evidence. Early work emphasized analytical competition and data-driven managerial thinking, arguing that

 Corresponding Author: rajabali.ghasempour@studenti.unicampania.it

 <https://doi.org/10.22105/masi.v3i4.111>



Licensee System Analytics. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0>).

firms could create advantage by embedding quantitative analysis in recurring decisions and business processes [1], [2]. Subsequent research shifted attention from isolated analytical models to organizational capabilities: data, technology, analytical talent, managerial commitment, and routines jointly determine whether analytics improves performance [3–5]. This capability view remains essential, but the object being governed is changing. Contemporary Artificial Intelligence (AI) systems no longer merely score alternatives or produce forecasts; they can generate explanations, synthesize unstructured information, converse with users, and increasingly interact with software tools. The central managerial question is therefore moving from “what does the model predict?” to “what should an AI-enabled system be allowed to decide and do?”

AI adoption in organizations has already challenged conventional distinctions between decision support and decision automation. Practical accounts emphasize that business value often emerges from targeted process redesign rather than ambitious technology-first projects [6]. Research on human–AI symbiosis similarly argues that computational systems and human decision makers possess different but potentially complementary strengths [7]. Organizational decision structures can therefore range from full human control to delegation, sequential hybrid arrangements, and aggregation of human and algorithmic judgments [8]. At a broader theoretical level, automation and augmentation are not mutually exclusive end states but interdependent organizational tendencies whose balance must be actively managed [9]. The management problem becomes more difficult as AI systems gain autonomy, learn from changing environments, and remain partly inscrutable [10].

This transition is accelerated by Large Language Models (LLMs) and agentic architectures. A conventional predictive model usually maps inputs to outputs within a predefined task. By contrast, an agentic system can decompose an objective into subgoals, retrieve information, select tools, execute steps, observe outcomes, and revise its plan. Such systems introduce a new layer between managerial intent and organizational action. They can potentially reduce coordination costs and accelerate decision cycles, but they also amplify familiar concerns surrounding opaque algorithmic control, accountability, and the redistribution of expertise [11], [12]. The resulting design problem cannot be solved by model accuracy alone. A system may be technically capable of performing an action yet organizationally unauthorized, insufficiently evidenced, poorly reversible, or socially unacceptable.

The motivation for this article is therefore a widening mismatch between the technical architecture of emerging AI agents and the governance architecture of organizations. Human–AI collaboration research has established the value of combining complementary capabilities [13], and team research has begun to conceptualize machines as active collaborators rather than passive tools [14]. Empirical studies of generative AI also show meaningful productivity effects in professional work, including faster completion and higher-quality outputs in writing tasks [15] and improved productivity among customer-support workers, especially less experienced employees [16]. Field evidence further suggests that performance gains depend on whether tasks fall inside or outside the uneven capability frontier of the technology [17]. These findings are consequential but do not yet provide a general management architecture for AI systems that can cross the boundary from recommendation to action.

This paper addresses that gap through a conceptual theory-building synthesis. It introduces Agentic Decision Intelligence (ADI) as a governance-bounded organizational capability in which AI systems do more than produce content or predictions: They participate in a monitored loop of sensing, evidence construction, reasoning, action, and learning. The innovation lies in treating autonomy as a conditional decision right rather than a fixed technical property. In this view, an AI agent's authority should expand or contract based on the stakes of a decision, uncertainty in the environment, reversibility of the proposed action, and the strength of accessible evidence. Governance is consequently embedded into the decision architecture instead of appended as an external compliance layer.

The paper makes four contributions. First, it differentiates analytical, generative, and agentic intelligence in management contexts and explains why the progression is organizational, not merely technological. Second, it proposes a five-layer ADI architecture that connects data and organizational knowledge to reasoning, tool-

mediated execution, and feedback. Third, it develops a decision-rights model supported by four controls, provenance, permissioning, escalation, and continuous assurance, that converts generic “human oversight” into operational design choices. Fourth, it formulates six research propositions and a staged implementation roadmap that can guide empirical work and organizational deployment. The main work is conceptual rather than empirical; no fabricated performance results are reported. The purpose is to offer a falsifiable, design-relevant framework that future studies can evaluate across domains such as operations, marketing, supply chains, finance, human resources, and public administration.

The remainder of the article is organized as follows. Section 2 synthesizes literature on analytics capabilities, human–AI collaboration, generative AI, autonomous agents, trust, explainability, and governance. Section 3 develops the ADI architecture. Section 4 specifies governance and decision rights. Section 5 presents research propositions and an implementation agenda. Section 6 concludes with implications, limitations, and future research directions.

2 | Literature Review and Conceptual Foundations

2.1 | From Business Analytics to AI-Enabled Decision Systems

Business analytics established the organizational logic on which ADI builds. Analytics creates value when firms develop not only technical infrastructure but also data quality, managerial skills, cross-functional routines, and mechanisms for translating insights into action [3–5]. This literature is important because it rejects technological determinism: Analytical capability is a configuration of resources and routines rather than a software acquisition. Agentic systems intensify this logic because the distance between model output and organizational action becomes shorter. When an AI system can execute as well as recommend, weaknesses in data governance, process ownership, or incentive alignment can propagate directly into operational consequences.

Traditional analytics can be organized around descriptive, predictive, and prescriptive functions. Descriptive analytics explains what has happened; predictive analytics estimates what is likely to happen; prescriptive analytics recommends what should be done. Agentic systems add a fourth operational dimension: controlled enactment. The system can decide how to pursue a delegated objective, invoke tools, coordinate subtasks, and revise execution based on observed feedback. This progression is summarized in *Table 1*. The distinction matters because the accountability burden changes as systems move from informational influence to behavioral intervention.

Table 1. Evolution from analytical support to ADI.

Main Control Question	Manager Role	System Role	Stage
Is the data valid and relevant?	Interpret evidence	Summarize patterns	Descriptive analytics
Is predictive performance adequate?	Judge predictions	Forecast or score	Predictive analytics
Are objectives and constraints correct?	Select or approve	Recommend actions	Prescriptive analytics
Is the output grounded and reliable?	Prompt, inspect, refine	Synthesize and explain	Generative AI
What may the system decide and execute?	Delegate and supervise	Plan, use tools, act, learn	ADI

As *Table 1* indicates, the defining change is not that AI becomes “smarter” in the abstract, but that organizational systems become capable of translating representations of managerial goals into sequences of consequential actions. This is precisely where established work on algorithmic management becomes relevant. Learning algorithms can reshape expertise, occupational boundaries, coordination, and control [11], while workplace algorithms create contested terrain around direction, evaluation, discipline, and worker autonomy

[12]. ADI must therefore be designed as an organizational control system, not simply an advanced interface to a model.

2.2 | Human–AI Complementarity, Delegation, and Trust

Research on human–AI complementarity provides the second conceptual foundation. Human decision makers are generally stronger in contextual interpretation, moral reasoning, handling ambiguous goals, and negotiating equivocal social situations, whereas AI systems can process large information spaces, apply consistent rules, and rapidly evaluate alternatives [7]. Hybrid intelligence research formalizes the idea that human and machine capabilities can be deliberately combined so that the joint system achieves outcomes neither component could reliably produce alone [13]. Related work on machines as teammates emphasizes that AI artifacts alter not only task allocation but also communication, coordination, and institutional arrangements [14].

However, complementarity does not imply that more autonomy is always better. Classic automation research distinguishes different types and levels of automation and shows that the optimal allocation of functions depends on reliability, consequences, and human-performance effects [27]. Trust research likewise demonstrates that effective automation requires calibrated reliance rather than maximal trust [26]. In AI contexts, trust is shaped by the system's representation, capabilities, performance, process transparency, and user expectations [28]. Under-trust wastes useful automation; over-trust creates automation bias and weakens human monitoring. The implication for agentic systems is that delegation should be dynamically bounded by evidence about competence and risk.

Generative AI strengthens this argument. Controlled experiments have shown substantial gains in speed and quality for some professional writing tasks [15], while workplace deployment in customer support has produced average productivity gains and disproportionately benefited less experienced workers [16]. Yet the “jagged technological frontier” observed in knowledge work means that seemingly similar tasks can differ sharply in whether current AI improves or degrades performance [17]. The governance architecture of ADI must therefore support selective delegation: autonomy should be high where competence is established and consequences are containable, but constrained where capability is uncertain or failure is costly.

2.3 | Large Language Models and the Emergence of Agentic Architectures

The third foundation concerns generative AI and autonomous agents. Large-scale foundation models create reusable capabilities across many downstream tasks, but they also concentrate risks because model limitations can propagate across applications [19]. The multidisciplinary literature on generative conversational AI highlights opportunities in productivity, creativity, decision support, and knowledge access alongside concerns regarding hallucination, bias, privacy, intellectual property, accountability, and misuse [18]. These characteristics make LLMs powerful cognitive interfaces but insufficient, by themselves, for reliable organizational agency.

Agentic architectures extend LLMs by adding memory, planning, tool use, environmental interaction, and feedback. Surveys of LLM-based autonomous agents converge on components such as perception, planning, memory, and action, while noting unresolved problems in evaluation, safety, long-horizon reliability, and social interaction [20], [21]. ReAct demonstrates how reasoning and acting can be interleaved so that a language model interacts with external environments rather than relying solely on internal generation [22]. Toolformer shows that language models can learn when and how to call external application programming interfaces, thereby combining general language capability with specialized tools [23]. Retrieval-Augmented Generation (RAG) connects parametric models to external knowledge stores, improving access to current or provenance-bearing information [24]. These techniques collectively enable a shift from conversational assistance toward action-oriented systems.

For management analytics, the technical significance of these advances is secondary to their organizational implication: An AI system can now traverse multiple decision stages that were previously separated by human

handoffs. It can retrieve a policy, inspect a dashboard, calculate an indicator, draft a recommendation, open a workflow, request approval, and write the resulting state back to a business system. This compression of the decision cycle creates value through speed and consistency but can also create tightly coupled failures. ADI therefore requires architectural separation between evidence, reasoning, authority, and execution.

2.4 | Explainability, Ethics, and Governance as Design Requirements

Human–AI interaction research emphasizes that systems should communicate what they can do, explain relevant behavior, support efficient correction, and adapt while preserving user control [25]. Explainability becomes especially important when AI outputs affect consequential decisions, although post-hoc explanations cannot substitute for inherently understandable models in every high-stakes setting [29]. Explainable Artificial Intelligence (XAI) research provides taxonomies of methods and identifies trade-offs between interpretability, performance, fidelity, and user needs [30]. For agentic systems, explanation must extend beyond why a prediction was generated to include what evidence was used, which tools were called, what actions were taken, and how the system's authority was determined.

Ethical and governance literatures reinforce this lifecycle perspective. AI ethics frameworks commonly converge on principles such as beneficence, non-maleficence, autonomy, justice, and explicability [31], while global guideline analyses identify recurring commitments to transparency, fairness, non-maleficence, responsibility, and privacy [32]. Earlier algorithm-ethics work also emphasizes opacity, bias, discrimination, autonomy, privacy, and responsibility as distinct but connected concerns [38]. The National Institute of Standards and Technology (NIST) Artificial Intelligence Risk Management Framework (AIRMF) operationalizes trustworthy AI around governance, mapping, measurement, and management [33], and its Generative Artificial Intelligence Profile extends the framework to generative-specific risks [34]. ISO/IEC 42001:2023 formalizes requirements for an Artificial Intelligence Management System (AIMS) [35]. The Organisation for Economic Co-operation and Development (OECD) AI principles, updated in 2024, emphasize trustworthy, human-centric AI [36], while the European Union AI Act establishes a risk-based regulatory regime for AI systems [37].

These sources point to a core conceptual conclusion: Governance cannot be reduced to an approval gate placed after technical development. Once AI systems can initiate actions, governance must be expressed as executable constraints within the architecture itself. Permissions, evidence requirements, human escalation, logging, and monitoring become part of the decision system's functional design. This conclusion motivates the ADI framework developed next.

3 | Agentic Decision Intelligence Framework

3.1 | Definition and Boundary Conditions

ADI is defined here as an organizational capability in which AI-enabled systems participate in a governance-bounded loop of sensing, evidence construction, reasoning, tool-mediated action, and learning in pursuit of delegated managerial objectives. Three elements distinguish ADI from ordinary generative assistance. First, the system is goal-directed: It receives an objective or decision mandate rather than only a prompt for content. Second, it is action-capable: It can invoke tools or workflows that alter an organizational state. Third, its autonomy is conditional: Authority is constrained by explicit policies, risk thresholds, evidence requirements, and escalation rules. ADI is therefore not synonymous with “autonomous AI.” It is a socio-technical configuration in which autonomy is granted selectively and remains auditable.

This definition has important boundary conditions. A chatbot that drafts a memo but cannot access authoritative organizational knowledge or initiate an action is generative assistance, not ADI. A robotic process automation script that executes deterministic rules without adaptive reasoning is automation, not ADI. A predictive model that triggers a pre-approved action may be algorithmic automation, but it becomes ADI only when it participates in a broader loop involving contextual evidence, adaptive planning, or

conditional decision rights. The framework is consequently oriented toward systems that combine heterogeneous reasoning and action capabilities under organizational control.

3.2|The Five-Layer Architecture

Fig. 1 presents the proposed architecture. The layers are deliberately separated because technical integration without governance separation encourages hidden coupling. Each layer has a distinct managerial purpose and produces artifacts that can be inspected independently.

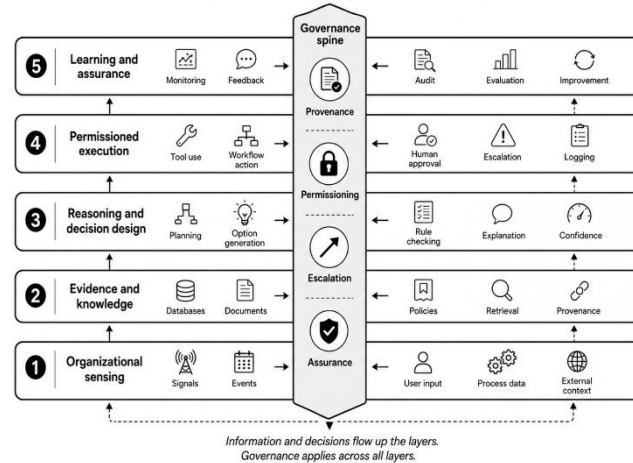


Fig. 1. Five-layer architecture of ADI.

Table 2. Functional layers of the proposed ADI architecture.

Required Control	Primary Failure Mode	Core Function	Layer
Input validation and scope checks	Incomplete or misleading context	Capture events, goals, states, and constraints	1. Organizational sensing
Provenance and freshness checks	Stale, biased, or ungrounded evidence	Retrieve and structure decision evidence	2. Evidence and knowledge
Structured reasoning and thresholds	Hallucinated logic or objective misalignment	Generate alternatives, evaluate constraints, plan	3. Reasoning and decision design
Least privilege, approvals, transaction limits	Unauthorized or irreversible action	Invoke tools and alter organizational state	4. Permissioned execution
Continuous monitoring and controlled updates	Silent drift or self-reinforcing error	Evaluate outcomes and adapt safely	5. Learning and assurance

Layer 1, organizational sensing, converts the environment into a decision context. Inputs may include structured enterprise data, events from operational systems, unstructured documents, user requests, customer communications, sensor streams, or managerial objectives. The central design requirement is not maximal data ingestion but contextual sufficiency. Agentic systems should know which objective they are pursuing, what constraints apply, which organizational unit owns the decision, and when the sensed state becomes stale. This layer prevents a common failure mode in which an apparently intelligent system reasons correctly about an incorrectly framed problem.

Layer 2 constructs an evidence base. Retrieval should distinguish authoritative internal sources, external information, model-generated hypotheses, and user-provided statements. RAG and related retrieval mechanisms are valuable because they can connect language-model reasoning to explicit knowledge stores [24], but retrieval alone does not create trustworthy evidence. Each evidence item should carry provenance, time, access status, and where possible, an estimate of reliability. The goal is to make factual grounding inspectable and to prevent generated content from being silently promoted into organizational fact.

Layer 3 performs reasoning and decision design. The system decomposes objectives, generates alternatives, checks constraints, and estimates expected consequences. Here, LLM reasoning may be combined with statistical models, optimization, simulation, business rules, or human judgment. This heterogeneous architecture is preferable to treating the LLM as a universal decision engine. In high-stakes decisions, interpretable models or explicit rules may be required even when a generative model is used to orchestrate the workflow [29], [30]. The output of Layer 3 should be a decision object: proposed action, supporting evidence, assumptions, confidence or uncertainty indicators, and the required level of authority.

Layer 4 is permissioned execution. It is the defining operational layer of ADI because recommendations become actions. Tool calls may create purchase orders, modify schedules, contact customers, update records, deploy software, allocate resources, or initiate approvals. Least-privilege design is essential: the agent should receive only the permissions required for the current task, with transaction limits and reversible operations where possible. The system must never infer authority solely from technical capability. Organizational policy, not model confidence, determines whether a tool may be used.

Layer 5 combines learning and assurance. Agentic systems require feedback on both task success and governance quality. A decision can achieve its immediate objective while violating process requirements, creating downstream harm, or weakening accountability. Evaluation should therefore include outcome metrics, process metrics, exception rates, human overrides, escalation frequency, unauthorized-attempt counts, evidence quality, and drift indicators. Learning from feedback should itself be governed: incident data can inform improved prompts, rules, retrieval, models, and policies, but self-modification without controlled validation would undermine assurance.

3.3 | The Closed-Loop Decision Cycle

The architecture operates as a closed loop rather than a linear pipeline. *Fig. 2* depicts the cycle. A trigger creates a decision episode; evidence is assembled; alternatives and plans are produced; a governance function determines allowable authority; the system either executes, requests human approval, or escalates; outcomes are observed; and the assurance layer updates performance records and, where appropriate, prompts organizational learning. The loop can be repeated until the objective is achieved, abandoned, or transferred to a human owner.

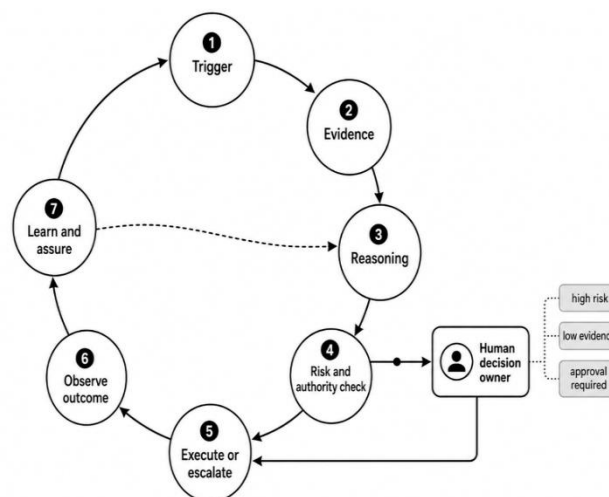


Fig. 2. Closed-loop decision cycle with dynamic authority allocation.

The cycle reframes human oversight. Rather than forcing a person to approve every action, which can create meaningless “rubber-stamp” supervision, human involvement is concentrated where judgment adds value. This principle follows the literature on calibrated trust and levels of automation [26–28]. Humans should intervene when uncertainty, novelty, value conflict, irreversible consequences, or authority boundaries exceed pre-defined tolerances. Routine, well-evidenced, reversible tasks can be delegated more extensively. Oversight thus becomes risk-sensitive and exception-oriented.

4 | Governance, Decision Rights, and Accountability

4.1 | A Four-Factor Decision-Rights Model

The central governance proposition of ADI is that autonomy should be allocated by decision characteristics rather than by application category alone. Four factors jointly determine an appropriate level of machine decision rights: Consequence severity, environmental uncertainty, action reversibility, and evidentiary strength. Consequence severity captures the magnitude and distribution of potential harm or loss. Environmental uncertainty captures how incomplete, unstable, or ambiguous the decision context is. Reversibility captures whether an erroneous action can be promptly and fully undone. Evidentiary strength captures the quality, recency, consistency, and provenance of information supporting the action.

These factors produce three practical authority zones. In the autonomous zone, consequences are limited, evidence is strong, uncertainty is low to moderate, and actions are reversible. The agent may execute within predefined limits. In the supervised zone, at least one risk factor is elevated; the system may prepare and partially execute the workflow but requires human authorization at a designated checkpoint. In the prohibited or specialist zone, consequences are severe, evidence is weak, uncertainty is high, or the action is effectively irreversible; the system may support analysis but lacks authority to execute. *Fig. 3* shows this concept as a governance envelope rather than a binary human-versus-machine choice.

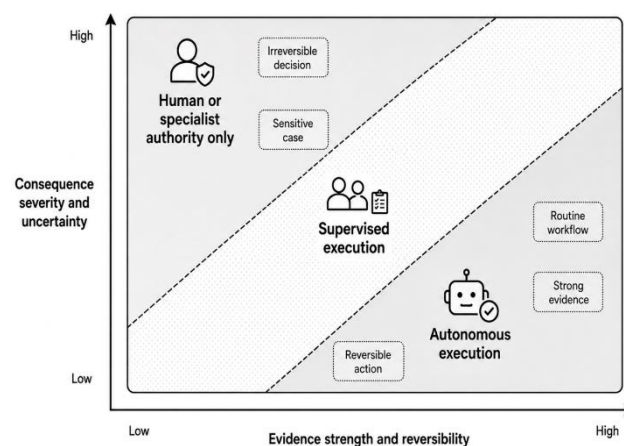


Fig. 3. Governance envelope for dynamic human–AI decision rights.

This model operationalizes a recurring theme in AI governance: Risk controls must be proportional to context. NIST frames AI risk management as a continuous process involving governance, mapping, measurement, and management [33], while NIST guidance for generative AI emphasizes lifecycle-specific risks and mitigations [34]. ISO/IEC 42001:2023 similarly treats AI governance as an organizational management system rather than a one-time technical assessment [35]. The ADI model translates these general principles into decision-level allocation of authority.

4.2 | Four Embedded Governance Mechanisms

The first mechanism is provenance. Every consequential decision episode should preserve a trace of the evidence used, model or tool versions involved, policy rules applied, prompts or structured instructions that materially affected the result, and human interventions. Provenance supports diagnosis, contestability, audit, and organizational learning. It also reduces a central danger of generative systems: Fluent outputs can obscure the distinction between retrieved fact, model inference, and fabricated content [18], [19], [40].

The second mechanism is permissioning. Agentic systems should follow least-privilege principles. Tool access should be scoped to role, task, data class, transaction value, and time. Permission boundaries must be enforced outside the language model whenever possible. A prompt stating “do not exceed \$10,000” is weaker than a

transactional control that technically prevents a larger commitment. Permissioning converts governance from language into infrastructure.

The third mechanism is escalation. Escalation should be triggered by explicit conditions such as conflicting evidence, low confidence, novel states, policy ambiguity, high-value transactions, sensitive data, unusual user requests, or failed validation checks. Effective escalation includes a clear human owner and a compact decision packet containing the proposed action, evidence, unresolved uncertainty, and reason for escalation. This is consistent with human–AI interaction guidance that systems should make their capabilities and limitations visible and facilitate user correction [25].

The fourth mechanism is continuous assurance. Pre-deployment testing is necessary but insufficient for systems that operate in changing organizational environments. Assurance should track technical performance, governance performance, and social consequences. Metrics may include outcome quality, override rates, escalation rates, incident severity, false-action frequency, retrieval provenance coverage, policy violations prevented, time-to-resolution, and differential performance across groups or contexts. The aim is calibrated autonomy: empirical evidence about performance can justify broader delegation, while incidents or drift should contract the authority envelope.

4.3 | Accountability and Organizational Role Design

Agentic systems complicate responsibility because action emerges from interactions among models, data, tools, policies, developers, process owners, and users. Ethics research has long identified traceability and responsibility as central problems of algorithmic mediation [31], [32], [38]. An organization should therefore avoid assigning accountability to “the AI.” Accountability must remain attached to human and institutional roles. At minimum, four roles should be explicit: the business decision owner, who owns objectives and consequences; the system owner, who owns technical operation; the risk or compliance owner, who defines control requirements; and the data/knowledge owner, who governs evidence sources. Smaller organizations may combine these roles, but the responsibilities should still be distinguishable.

This role design also addresses the automation–augmentation paradox [9]. If management treats AI exclusively as automation, human expertise can atrophy and responsibility may be displaced. If management treats AI exclusively as augmentation, organizations may retain costly manual checkpoints that add little value. ADI instead treats human involvement as a portfolio of decision rights: Some decisions are delegated, some are jointly produced, some require approval, and some remain exclusively human. The portfolio should evolve with evidence about system capability and organizational risk appetite.

Table 3. Governance control matrix for ADI.

Example	Mandatory Controls	Default Authority	Risk Condition
Reschedule a low-priority internal meeting	Logging, policy check, transaction cap	Autonomous within limits	Low consequence; strong evidence; reversible
Change a replenishment quantity above normal band	Evidence packet, approval gate, rollback	Supervised execution	Moderate consequence or uncertainty
Commit substantial capital expenditure	Independent validation, dual approval, audit trail	Human authorization	High consequence or low reversibility
Respond to an unusual compliance case	Escalation, extra evidence, specialist review	No autonomous execution	High uncertainty or conflicting evidence
Action outside delegated legal authority	Technical denial, monitoring, incident alert	Analysis only or no access	Prohibited policy domain

Table 3 illustrates that governance should be implemented as a combination of technical controls and human role assignments. The same underlying AI model may therefore operate with different autonomy in different processes. This context sensitivity is preferable to blanket claims that a model is “safe” or “unsafe.”

5 | Research Propositions and Implementation Agenda

5.1 | Research Propositions

The framework generates six propositions intended for empirical testing. They are stated at a level that can support survey research, field experiments, process-mining studies, design-science evaluations, or longitudinal case research.

Proposition 1. The positive effect of agentic AI capability on decision-cycle performance is stronger when organizational knowledge is provenance-rich and operationally retrievable than when knowledge is fragmented or weakly governed.

Proposition 2. Conditional autonomy based on consequence severity, uncertainty, reversibility, and evidentiary strength produces higher joint decision quality and lower incident severity than either uniform human approval or uniform machine autonomy.

Proposition 3. Provenance visibility improves calibrated human reliance on agentic recommendations by enabling users to distinguish authoritative evidence, model inference, and uncertainty.

Proposition 4. Permission controls enforced outside the generative model reduce the relationship between model error and organizational harm more effectively than instruction-only constraints embedded in prompts.

Proposition 5. Exception-oriented human oversight produces higher supervisory effectiveness than universal approval when routine decision volume is high and escalation criteria are well specified.

Proposition 6. Continuous assurance that can contract or expand an agent's authority based on observed performance increases long-run organizational value relative to static pre-deployment approval.

Proposition 1 connects analytics-capability research [3–5] to agentic architectures by arguing that knowledge governance becomes a direct antecedent of action quality. *Propositions 2* and *5* extend classic automation and trust theory [26–28] into dynamic decision-rights allocation. *Proposition 3* connects explainability to reliance calibration rather than treating explanation as an intrinsic good [29], [30]. *Proposition 4* differentiates policy statements from enforceable infrastructure, an increasingly important distinction when language models can plan around loosely phrased constraints. *Proposition 6* treats autonomy as an adaptive organizational resource whose scope should depend on evidence accumulated through use.

5.2 | Staged Organizational Implementation

Organizations should not implement ADI by moving directly from a chatbot pilot to broad autonomous execution. A staged roadmap reduces both technical and institutional risk. *Stage 1* is assistive analytics: the system retrieves information, summarizes evidence, and drafts options, while humans retain all decision rights. *Stage 2* is supervised agency: The system can construct plans and prepare tool calls, but each consequential action requires explicit approval. *Stage 3* is bounded autonomy: The system executes low-risk, reversible actions within defined limits and escalates exceptions. *Stage 4* is adaptive autonomy: The authority envelope is adjusted through continuous assurance, performance evidence, and formal governance review.

The transition between stages should be evidence-based. Useful readiness criteria include retrieval provenance coverage, test-set performance on representative cases, exception-handling accuracy, frequency and causes of human overrides, rollback success, incident severity, policy compliance, and user understanding of system limitations. Productivity alone is not a sufficient promotion criterion. A system that saves time while increasing tail risk or degrading accountability should not receive broader decision rights.

Implementation also requires process redesign. Organizations should identify decision episodes rather than merely automate tasks. A decision episode has an owner, trigger, objective, information requirements, permissible actions, consequences, and escalation path. Mapping these elements often reveals that existing processes rely on tacit judgment or undocumented exceptions. ADI deployment can therefore expose

governance debt that predates AI. The appropriate response is not to encode ambiguity into prompts but to clarify policies, decision rights, and exception ownership.

5.3 | Evaluation Metrics for Management Analytics

Agentic systems require a broader performance model than predictive analytics. Four metric families are proposed. Decision quality includes accuracy, economic value, service level, or domain-specific outcomes. Process efficiency includes cycle time, cost, throughput, and coordination burden. Governance quality includes provenance coverage, policy compliance, escalation appropriateness, unauthorized-action prevention, and audit completeness. Human-system quality includes calibrated reliance, override quality, perceived control, workload, learning, and the preservation of relevant expertise. Evaluating all four families avoids the narrow optimization of speed at the expense of accountability or resilience.

A particularly important metric is recoverability. Because no complex system will be error-free, organizations should measure how quickly an erroneous action is detected, contained, reversed, and learned from. Reversibility is not only a decision-rights factor but also an engineering target. Workflows can be deliberately designed with staged commits, sandbox environments, transaction caps, and compensating actions so that AI errors remain locally containable. This design-for-recovery perspective is likely to become a defining capability of mature agentic organizations.

5.4 | Implications for Management Theory and Practice

For management theory, ADI suggests that the unit of analysis should shift from the “AI tool” toward the human–AI decision system. Performance emerges from the configuration of data, models, knowledge, authority, users, workflows, and controls. This perspective aligns with research showing that AI management involves interdependent choices around autonomy, learning, and inscrutability [10]. It also offers a bridge between business analytics research, which emphasizes capabilities and value creation, and organizational AI research, which emphasizes control, expertise, and institutional consequences.

For practice, the framework provides a concrete design question: what is the smallest safe authority that allows the system to create meaningful value? Starting with bounded permissions and expanding them through evidence is more defensible than granting broad access and attempting to monitor after the fact. This principle is compatible with contemporary governance frameworks and standards [33–37], yet it is expressed in operational terms that can be implemented by process owners, AI teams, security teams, and managers.

6 | Conclusion

The evolution of management analytics is entering a qualitatively different stage. Predictive and prescriptive systems primarily influenced decisions by producing information or recommendations; emerging agentic systems can participate directly in the sequence that turns managerial intent into organizational action. This change creates substantial opportunities for faster decision cycles, lower coordination costs, scalable expertise, and more adaptive operations, but it also changes the locus of risk. Errors are no longer confined to a dashboard or generated paragraph; they can propagate through tools, workflows, and transactions.

This article developed the concept of ADI as a governance-bounded capability that integrates sensing, evidence, reasoning, permissioned execution, and learning. Its central argument is that autonomy should be treated as a conditional decision right. Consequence severity, environmental uncertainty, reversibility, and evidentiary strength determine whether a system may act autonomously, operate under supervision, or remain limited to analysis. Provenance, permissioning, escalation, and continuous assurance embed governance into the architecture itself. The framework thereby moves beyond generic calls for “human oversight” and specifies where human authority is most valuable.

The paper also contributes six propositions and a staged implementation agenda. These propositions invite empirical investigation of how knowledge provenance, dynamic delegation, explainability, permission controls, exception-oriented oversight, and adaptive assurance affect organizational outcomes. For managers,

the framework encourages a disciplined progression from assistive use to supervised agency, bounded autonomy, and eventually adaptive autonomy only when evidence justifies broader delegation. The resulting management objective is not maximal automation but accountable decision performance.

Several limitations should guide future work. The framework is conceptual and has not yet been validated with primary organizational data; its propositions should therefore be tested rather than treated as established causal claims. The four decision-rights factors may require domain-specific weighting, and some sectors will impose legal or ethical constraints that override any general autonomy logic. The architecture also abstracts from strategic behavior by users, adversarial attacks, multi-agent coordination, and long-term changes in workforce expertise. Future research should develop validated measurement scales for ADI capability, conduct field experiments comparing static and dynamic oversight regimes, examine multi-agent governance, quantify recoverability and tail risk, and study how autonomous decision systems reshape organizational structure, professional identity, and accountability over time. These directions are essential if management analytics is to progress from intelligent recommendations to responsible, governable action.

Acknowledgments

The authors thank colleagues and reviewers for providing constructive comments on future versions of this conceptual work. No external data collection or participant involvement was undertaken for this article.

Author Contribution

All authors contributed significantly to this research. All authors have read and approved the submitted version of the manuscript.

Funding

This research received no external funding.

Data Availability

No new empirical data were generated or analyzed in this conceptual study. All sources used to develop the framework are publicly identifiable in the reference list.

Conflicts of Interest

The authors declare no conflict of interest.

References

- [1] Davenport, T. H., & Harris, J. G. (2007). *Competing on analytics: The new science of winning*. Harvard Business School Press. <https://www.amazon.com/s?k=9781422103326&i=stripbooks&linkCode=qs>
- [2] Provost, F., & Fawcett, T. (2013). *Data science for business: What you need to know about data mining and data-analytic thinking*. O'Reilly Media. <https://www.amazon.com/Data-Science-Business-Data-Analytic-Thinking/dp/1449361323>
- [3] Gupta, M., & George, J. F. (2016). Toward the development of a big data analytics capability. *Information & Management*, 53(8), 1049–1064. <https://doi.org/10.1016/j.im.2016.07.004>
- [4] Wamba, S. F., Gunasekaran, A., Akter, S., Ren, S. J. f., Dubey, R., & Childe, S. J. (2017). Big data analytics and firm performance: Effects of dynamic capabilities. *Journal of Business Research*, 70, 356–365. <https://doi.org/10.1016/j.jbusres.2016.08.009>
- [5] Mikalef, P., Krogstie, J., Pappas, I. O., & Pavlou, P. (2020). Exploring the relationship between big data analytics capability and competitive performance: The mediating roles of dynamic and operational capabilities. *Information & Management*, 57(2), 103169. <https://doi.org/10.1016/j.im.2019.05.004>

- [6] Davenport, T. H., & Ronanki, R. (2018). Artificial intelligence for the real world. *Harvard Business Review*, 96(1), 108–116.
https://academicjournals.org/asset/documents/Artificial_Intelligence_for_the_Real_World_-_HBR.pdf
- [7] Jarrahi, M. H. (2018). Artificial intelligence and the future of work: Human–AI symbiosis in organizational decision making. *Business Horizons*, 61(4), 577–586. <https://doi.org/10.1016/j.bushor.2018.03.007>
- [8] Shrestha, Y. R., Ben-Menahem, S. M., & von Krogh, G. (2019). Organizational decision-making structures in the age of artificial intelligence. *California Management Review*, 61(4), 66–83.
<https://doi.org/10.1177/0008125619862257>
- [9] 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>
- [10] Berente, N., Gu, B., Recker, J., & Santhanam, R. (2021). Managing artificial intelligence. *MIS Quarterly*, 45(3), 1433–1450. <https://doi.org/10.25300/MISQ/2021/16274>
- [11] Faraj, S., Pachidi, S., & Sayegh, K. (2018). Working and organizing in the age of the learning algorithm. *Information and Organization*, 28(1), 62–70. <https://doi.org/10.1016/j.infoandorg.2018.02.005>
- [12] Kellogg, K. C., Valentine, M. A., & Christin, A. (2020). Algorithms at work: The new contested terrain of control. *Academy of Management Annals*, 14(1), 366–410. <https://doi.org/10.5465/annals.2018.0174>
- [13] Dellermann, D., Ebel, P., Söllner, M., & Leimeister, J. M. (2019). Hybrid intelligence. *Business & Information Systems Engineering*, 61(5), 637–643. <https://doi.org/10.1007/s12599-019-00595-2>
- [14] Seeber, I., Bittner, E., Briggs, R. O., De Vreede, T., De Vreede, G. J., Elkins, A., Maier, R., Merz, A. B., Oeste-Reiß, S., Randrup, N., Schwabe, G., & Söllner, M. (2020). Machines as teammates: A research agenda on AI in team collaboration. *Information & Management*, 57(2), 103174.
<https://doi.org/10.1016/j.im.2019.103174>
- [15] Noy, S., & Zhang, W. (2023). Experimental evidence on the productivity effects of generative artificial intelligence. *Science*, 381(6654), 187–192. <https://doi.org/10.1126/science.adh2586>
- [16] Brynjolfsson, E., Li, D., & Raymond, L. R. (2025). Generative AI at work. *The Quarterly Journal of Economics*, 140(2), 889–942. <https://doi.org/10.1093/qje/qjae044>
- [17] Dell'Acqua, F., McFowland III, E., Mollick, E., Lifshitz-Assaf, H., Kellogg, K. C., Rajendran, S., ... & Lakhani, K. R. (2023). *Navigating the jagged technological frontier: Field experimental evidence of the effects of AI on knowledge worker productivity and quality*. <https://doi.org/10.2139/ssrn.4573321>
- [18] Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., Baabdullah, A. M., Koohang, A., Raghavan, V., Ahuja, M., et al. (2023). “So what if ChatGPT wrote it?” Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*, 71, 102642.
<https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- [19] Bommasani, R., Hudson, D. A., Adeli, E., Altman, R., Arora, S., von Arx, S., Bernstein, M. S., Bohg, J., Bosselut, A., Brunskill, E., et al. (2021). *On the opportunities and risks of foundation models*.
<https://doi.org/10.48550/arXiv.2108.07258>
- [20] Wang, L., Ma, C., Feng, X., Zhang, Z., Yang, H., Zhang, J., Chen, Z., Tang, J., Chen, X., Lin, Y., Zhao, W. X., Wei, Z., & Wen, J. (2024). A survey on large language model based autonomous agents. *Frontiers of Computer Science*, 18, 186345. <https://doi.org/10.1007/s11704-024-40231-1>
- [21] Xi, Z., Chen, W., Guo, X., He, W., Ding, Y., Hong, B., Zhang, M., Wang, J., Jin, S., Zhou, E., et al. (2025). The rise and potential of large language model based agents: A survey. *Science China Information Sciences*, 68, 121101. <https://doi.org/10.1007/s11432-024-4222-0>
- [22] Yao, S., Zhao, J., Yu, D., Du, N., Shafran, I., Narasimhan, K. R., & Cao, Y. (2023). ReAct: Synergizing reasoning and acting in language models. <https://doi.org/10.48550/arXiv.2210.03629>
- [23] Schick, T., Dwivedi-Yu, J., Dessì, R., Raileanu, R., Lomeli, M., Hambro, E., ... & Scialom, T. (2023). Toolformer: Language models can teach themselves to use tools. *Advances in Neural Information Processing Systems*, 36, 68539–68551. <https://doi.org/10.52202/075280-2997>
- [24] Lewis, P., Perez, E., Piktus, A., Petroni, F., Karpukhin, V., Goyal, N., Küttler, H., Lewis, M., Yih, W. t., Rocktäschel, T., Riedel, S., & Kiela, D. (2020). Retrieval-augmented generation for knowledge-intensive

- NLP tasks. *Advances in Neural Information Processing Systems*, 33, 9459–9474.
<https://doi.org/10.5555/3495724.3496517>
- [25] Amershi, S., Weld, D., Vorvoreanu, M., Fournery, A., Nushi, B., Collisson, P., & Horvitz, E. (2019). Guidelines for human-AI interaction. In *Proceedings of the 2019 Chi Conference on Human Factors in Computing Systems* (pp. 1-13). Association for Computing Machinery.
<https://doi.org/10.1145/3290605.3300233>
- [26] Lee, J. D., & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46(1), 50–80. https://doi.org/10.1518/hfes.46.1.50_30392
- [27] Parasuraman, R., Sheridan, T. B., & Wickens, C. D. (2000). A model for types and levels of human interaction with automation. *IEEE Transactions on Systems, Man, and Cybernetics—Part A: Systems and Humans*, 30(3), 286–297. <https://doi.org/10.1109/3468.844354>
- [28] Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of Management Annals*, 14(2), 627–660. <https://doi.org/10.5465/annals.2018.0057>
- [29] Rudin, C. (2019). Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead. *Nature Machine Intelligence*, 1, 206–215. <https://doi.org/10.1038/s42256-019-0048-x>
- [30] Barredo Arrieta, A., Díaz-Rodríguez, N., Del Ser, J., Bennetot, A., Tabik, S., Barbado, A., García, S., Gil-Lopez, S., Molina, D., Benjamins, R., Chatila, R., & Herrera, F. (2020). Explainable artificial intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information Fusion*, 58, 82–115. <https://doi.org/10.1016/j.inffus.2019.12.012>
- [31] Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., & Vayena, E. (2018). AI4People—An ethical framework for a good AI society: Opportunities, risks, principles, and recommendations. *Minds and Machines*, 28(4), 689–707. <https://doi.org/10.1007/s11023-018-9482-5>
- [32] Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1, 389–399. <https://doi.org/10.1038/s42256-019-0088-2>
- [33] Tabassi, E. (2023). *Artificial intelligence risk management framework (AI RMF 1.0) (NIST AI 100-1)*.
<https://doi.org/10.6028/NIST.AI.100-1>
- [34] Autio, C., Schwartz, R., Dunietz, J., Jain, S., Stanley, M., Tabassi, E., Hall, P., & Roberts, K. (2024). *Artificial intelligence risk management framework: Generative artificial intelligence profile (NIST AI 600-1)*.
<https://doi.org/10.6028/NIST.AI.600-1>
- [35] International Organization for Standardization. (2023). *ISO/IEC 42001:2023 information technology—artificial intelligence—management system*. <https://www.onetrust.com/resources/navigating-the-iso-42001-framework-ebook/>
- [36] Organisation for Economic Co-operation and Development. (2024). *OECD AI principles*.
<https://oecd.ai/en/ai-principles>
- [37] European Parliament & Council of the European Union. (2024). *Regulation (EU) 2024/1689 of 13 June 2024 laying down harmonised rules on artificial intelligence (artificial intelligence act)*.
<https://www.wipo.int/wipolex/en/legislation/details/24023>
- [38] Mittelstadt, B. D., Allo, P., Taddeo, M., Wachter, S., & Floridi, L. (2016). The ethics of algorithms: Mapping the debate. *Big Data & Society*, 3(2). <https://doi.org/10.1177/2053951716679679>
- [39] Doshi-Velez, F., & Kim, B. (2017). *Towards a rigorous science of interpretable machine learning*.
<https://doi.org/10.48550/arXiv.1702.08608>
- [40] Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021, March). On the dangers of stochastic parrots: Can language models be too big?. *Proceedings of the 2021 ACM Conference on fairness, Accountability, and Transparency* (pp. 610-623). Association for Computing Machinery.
<https://doi.org/10.1145/3442188.3445922>