



Authority and Responsibility Allocation in Human–AI Collaborative Decision-Making: Governance Mechanisms for Enterprises

Fang Sun^{1,*}

¹Department of Business Administration, Dongshin University, Jeollanam-do 58245, Republic of Korea

Abstract

Artificial intelligence is moving from analytical support toward active participation in enterprise decisions, creating an organization-design problem: firms must decide which rights may be delegated to AI and how responsibility should follow the actors who can prevent, challenge, or remedy failure. Existing work explains automation, augmentation, delegation, human oversight, and responsible AI governance, but does not reveal how specific transfers of decision authority create responsibility gaps inside a focal enterprise decision. This conceptual paper develops a contingency governance framework through a transparent theory-synthesis procedure. A purposive corpus of 44 peer-reviewed studies, standards, and regulatory sources was assembled through anchor studies, targeted keyword searches, and citation chaining. First-order authority and responsibility terms were coded, compared, and abstracted until two successive search iterations produced no new categories. The resulting framework

distinguishes seven decision rights—information access, recommendation, selection, approval, veto, execution, and escalation—and five responsibility domains—system design, decision process, outcome stewardship, oversight, and remediation. Its central mechanism is rights-control-responsibility alignment: delegating a right shifts effective control and evidence access, while governance fails when the responsible actor lacks the competence, authority, or information to intervene. Decision exposure and AI autonomy determine four governance archetypes, while AI reliability conditions the permissible scope of selection and execution rights. Eight empirically testable propositions specify mechanisms, moderators, competing explanations, and falsification conditions. Two worked applications show how the architecture produces more precise governance than a generic human-in-the-loop requirement. The paper contributes a decision-level theory of enterprise AI governance and provides managers with an auditable method for allocating rights, responsibilities, evidence, and lifecycle controls.

Keywords: human-AI collaboration, enterprise decision-making, decision rights, responsibility allocation, AI governance, meaningful human control.



Submitted: 12 June 2026
Revised: 21 July 2026
Accepted: 28 July 2026
Published: 09 August 2026

Vol. 2, No. 3, 2026.
 10.62762/TSSR.2026.930520

*Corresponding author:
✉ Fang Sun
sf2025@dsu.ac.kr

Citation

Sun, F. (2026). Authority and Responsibility Allocation in Human–AI Collaborative Decision-Making: Governance Mechanisms for Enterprises. *ICCK Transactions on Systems Safety and Reliability*, 2(3), 192–206.

© 2026 ICCK (Institute of Central Computation and Knowledge)

1 Introduction

Artificial intelligence (AI) no longer enters organizations only as a back-office analytical tool. It increasingly generates options, ranks alternatives, recommends actions and, in some settings, executes decisions with limited human intervention. This development intensifies the automation–augmentation paradox: the same technology can substitute for human judgment while simultaneously increasing the need for human interpretation, coordination and oversight [1]. Managing AI therefore requires more than selecting a technically capable model. It requires organizations to redesign structures, routines and governance arrangements around a new class of agentic technology [2]. Generative AI extends this challenge because its outputs are adaptive, probabilistic and communicative rather than confined to stable prediction tasks, which means that the decision rights governing what AI may recommend, select, or execute cannot be fixed at deployment but must be revisited as the system’s output domain expands [3].

The managerial significance of AI is commonly framed in terms of productivity, business value or adoption. Such perspectives are necessary but incomplete. Reviews of AI and business value show that organizational outcomes depend on complementary resources, managerial capabilities and process redesign rather than on the technology alone, implying that governance arrangements—including how decision rights are distributed and responsibility assigned—are as consequential as the AI system itself [4]. More importantly, evidence on human-AI collaboration does not support a general presumption that a combined team will outperform either party. A large meta-analysis finds that human-AI combinations often perform below the single best performer—human or AI—especially in decision tasks [5]. The relevant managerial question is therefore not whether collaboration is inherently superior, but how an enterprise should configure authority, control and review so that human and machine capabilities are combined under appropriate conditions.

Existing frameworks illuminate important parts of this problem but leave a specific theoretical gap. Automation-augmentation research distinguishes broad modes of human-AI work [1]. Organization-design and agentic information-systems research show that tasks, information and authority can be delegated in different directions [6, 7].

Responsible AI governance and lifecycle research specify organizational practices, monitoring and auditing [8, 9]. Meaningful-human-control research explains why nominal oversight may be ineffective [10, 11]. These perspectives do not, however, identify which discrete right is transferred at a particular decision stage, which actor possesses effective control after the transfer, and which responsibility domain should follow that control. As a result, they cannot fully explain why a process may contain a human reviewer and still produce symbolic approval, delayed intervention or post-incident responsibility shifting.

The paper addresses this limitation through a rights-control-responsibility alignment mechanism. A transfer of recommendation, selection, approval, veto, execution or escalation authority changes who can shape the decision and who can access relevant evidence. Responsibility becomes operationally meaningful only when the responsible actor has the competence, information and authority needed to prevent, detect, contest or remedy the relevant failure. Misalignment arises when formal responsibility remains with an actor whose effective control has moved elsewhere, or when an actor exercises substantial control without corresponding responsibility. The seven-rights and five-responsibility architecture makes these hidden structural gaps observable, comparable and governable.

A second unresolved issue concerns AI reliability. Reliability is mentioned in the literature through accuracy, calibration, robustness, stability, drift and monitoring, but it is often conflated with general capability or treated as an outcome of good governance. This paper treats AI reliability as an antecedent and dynamic contingency of right allocation: lower or uncertain reliability narrows permissible selection and execution rights and increases human approval, escalation and monitoring. Governance does not itself create reliable performance; it generates the evidence and feedback required to assess, sustain and revise warranted confidence in reliability [12, 13].

This paper asks three questions. First, how should decision rights be distributed between human managers and AI systems? Second, how should responsibility be allocated among business owners, technical teams, executives and external providers? Third, how do decision exposure, AI autonomy and AI reliability shape the governance configuration? The paper contributes by specifying a

decision-level explanatory mechanism, deriving seven rights and five responsibility domains through an auditable theory-building procedure, distinguishing normative design principles from empirically testable propositions, and demonstrating the framework through low- and high-exposure applications.

2 Transparent Theory-Building Procedure

2.1 Conceptual design and unit of analysis

The study is a theory-synthesis and model-building conceptual article rather than an empirical study or a systematic review. Conceptual research requires an explicit design that explains how theoretical materials were selected, combined and transformed into a new explanation [14]. The focal unit is an enterprise decision in which an AI system materially shapes information generation, option evaluation, choice, approval, implementation or exception handling. This unit is narrower than general AI adoption but broader than AI-assisted recommendation because the system may receive operative rights and may affect responsibility before, during and after the focal choice.

The literature boundaries were defined by the theoretical problem. Four streams were included: human-AI collaboration and delegation; organization design, work design and algorithmic management; calibrated reliance, explainability and meaningful human control; and responsible AI governance, auditing and regulation. Studies on the allocation of decision rights and responsibility show how authority over choices must be explicitly distributed between human and AI actors, and how responsibility gaps arise when that distribution leaves accountable actors without effective control [15, 16]. Algorithmic management research shows how organizational objectives become embedded in data capture, monitoring and automated intervention [17]. Purely technical studies were included only when they clarified a governance-relevant property such as reliability, drift, robustness or logging. Studies were excluded when they addressed model performance without organizational implications, consumer adoption without decision authority, or ethical principles without an implementable governance construct.

2.2 Literature identification and selection

The corpus was assembled in July 2026 through a transparent seed-and-snowballing protocol. Anchor works were selected to represent the main theoretical streams: the automation-augmentation paradox [1],

human-AI decision making as an organization-design problem [6], delegation to and from agentic information systems [7], responsible AI governance reviews [8, 21], lifecycle decision systems [9], strategic organizational decision support [18], and multi-level meaningful human control [11]. Backward and forward citation chaining was combined with targeted searches on publisher platforms using query families such as "human-AI decision authority," "AI delegation approval veto execution," "organizational AI governance responsibility," "meaningful human control escalation," and "AI reliability decision rights."

Sources were retained when they met at least one of four criteria: they defined a distinct authority relation; specified a responsibility, accountability or remediation domain; identified a contextual condition that changes permissible delegation; or described a control mechanism connecting evidence to intervention. Priority was given to peer-reviewed English-language work published from 2020 to July 2026, while authoritative standards and regulation were included because they operationalize lifecycle controls. The final corpus comprised 44 sources: 40 peer-reviewed articles or proceedings, two standards or regulatory instruments, one authoritative risk-management framework and one conceptual-method source.

2.3 Concept extraction, abstraction and category derivation

A concept matrix was used to record the unit of analysis, authority terms, responsibility terms, contextual conditions, control mechanisms and proposed outcomes for each source. First-order coding retained the vocabulary used by the literature, including access, advise, recommend, rank, choose, approve, override, stop, implement, monitor, audit, compensate and learn. Constant comparison was then used to distinguish terms that appeared similar but changed different organizational relations. Recommendation was separated from selection because proposing an option does not determine the chosen action. Approval was separated from veto because permission to activate a decision differs from authority to suspend or reject it. Execution was separated from selection because implementation may be automated after a human choice, or a system may select within a domain while a human executes.

Responsibility terms were abstracted by the object and timing of the duty. Design responsibility concerns the problem definition, data, model

Table 1. Comparison with prior human-AI and AI governance frameworks.

Framework or stream	Primary unit	Authority treatment	Responsibility treatment	Unresolved problem / added contribution
Automation-augmentation paradox [1]	Human-AI relationship task	Broad modes of substitution and complementarity	Not decomposed across the decision lifecycle	Cannot locate the specific right whose transfer creates a control gap
Organization design [6]	Tasks, information and authority	Authority recognized as a design variable	Responsibility remains broad	Does not provide an operational rights or responsibility portfolio
Agentic IS delegation [7]	Delegation to and from agentic artifacts	Direction and scope of delegation	Limited treatment of lifecycle responsibility	Cannot connect delegated action to design, outcome and remediation duties
Strategic human-AI decision support [18]	Organizational decision process under uncertainty	Distributes tasks and support across stages	Emphasizes continuing human responsibility	Lacks granular rights, named responsibility domains and falsifiable alignment mechanisms
Organizational governance [8, 19–23]	AI Organization-level structures and practices	Roles, policies and oversight at a broad level	Strong lifecycle and accountability emphasis	Does not specify who may recommend, select, approve, veto, execute or escalate a focal decision
Lifecycle and auditing [9, 24–26]	System stages, evidence and assurance	Authority is secondary	Strong monitoring, audit and follow-up logic	Does not explain how authority allocation produces responsibility gaps
Meaningful control [10, 11]	human Human oversight and control capability	Emphasizes override, intervention and feedback	Human responsibility is central	Does not integrate a complete rights portfolio with differentiated responsibility
Proposed framework	Focal enterprise decision	Seven separable decision rights	Five lifecycle responsibility domains	Explains how transfers of rights alter effective control and create observable alignment or misalignment

and validation. Decision-process responsibility concerns approved use and contemporaneous review. Outcome stewardship concerns cumulative effects after individual decisions. Oversight responsibility concerns the governance system itself. Remediation concerns correction, redress and organizational learning. Candidate categories were retained only when they were conceptually non-redundant, had a distinct governance consequence and could be linked to observable evidence.

2.4 Saturation, negative cases and traceability

Conceptual saturation was assessed at the category level rather than by counting publications. Search and coding proceeded iteratively. Two successive targeted search iterations produced additional examples and supporting evidence but did not yield a new decision right or responsibility domain. Negative cases tested the boundaries: protected non-reliance was mapped to veto or override rather than treated as an eighth right, while access to human reconsideration was mapped to escalation, contestability and remediation. The derivation of the seven rights and five responsibility domains is documented in Sections 2.3 and 4, while Table 2 provides a proposition-level audit trail linking each empirically testable proposition to its theoretical foundation, causal mechanism, observable outcome and rejection condition.

A structured comparison of the proposed framework with prior human-AI and AI governance frameworks is presented in Table 1, which highlights the specific unresolved problem that each stream leaves open and the added contribution of the present framework.

2.5 Normative design logic and empirical propositions

The framework separates two forms of theoretical claim. Sections 4 and 5 present normative design principles: they specify how rights and responsibilities should be configured if the enterprise seeks governability, meaningful control and auditable responsibility. Section 6 presents empirically testable propositions about organizational behavior and outcomes. Each proposition identifies a dependent variable, causal mechanism, boundary condition, competing explanation and observable rejection condition. This distinction avoids treating recommended governance as a prediction of what enterprises necessarily do.

3 Theoretical Foundations

3.1 Human-AI collaboration as differentiated delegation

Human-AI collaboration covers configurations that should not be conflated. In automation,

Table 2. Theoretical foundations and empirical testability of the propositions

Proposition	Theoretical foundation	Causal mechanism	Observable outcome / dependent variable	Boundary and rejection condition
Proposition 1	Reliability, bounded domains and reversibility [5, 12, 27]	Reliability uncertainty raises expected autonomous-error cost	Scope of AI selection and execution rights; approval and escalation intensity	Reversibility, error propagation and cost pressure; rejected if rights expand as reliability declines after controls.
Proposition 2	Exposure, legitimacy and human authority [10, 28, 29]	Exposure raises legitimacy and accountability costs	Human retained rights and governance intensity	Ethical salience, stakeholder reach and regulation; rejected if exposure rises without stronger retained rights or controls.
Proposition 3	Complementarity and information asymmetry [5, 27]	Capability-right fit integrates complementary information	Joint decision quality and preventable error	Management quality, task novelty and information asymmetry; rejected if status-based allocation performs equally or better.
Proposition 4	Symbolic oversight and control deficits [10, 22, 23]	Control deficits create perceived futility and information gaps	Symbolic approval, intervention delay and blame shifting	Workload, vendor opacity and ethical culture; rejected if severe mismatch coexists with effective challenge and rapid intervention.
Proposition 5	Appropriate reliance and explanation [30, 31]	Diagnostic explanations improve case-level discrimination	Discrimination accuracy and calibrated reliance	User expertise, complexity and cognitive forcing; rejected if acceptance rises without improved discrimination.
Proposition 6	Vendor opacity and audit access [8, 19–23, 32, 33]	Opacity reduces evidence access and governability	Contractual controls and permissible operative rights	Audit access, assurance and switching costs; rejected if opaque providers receive broader rights without stronger controls.
Proposition 7	Continuous auditing and lifecycle traceability [9, 24–26]	Traceability localizes drift or failure sources	Speed and accuracy of rights recalibration	Log completeness, suspension authority and governance maturity; rejected if monitoring completeness has no relation to recalibration.
Proposition 8	Contestability, redress and stakeholder legitimacy [25, 29, 34]	External challenge reveals hidden errors and procedural harms	Error discovery, correction speed, legitimacy and learning	Stakeholder reach, review independence and accessibility; rejected if contestability produces no improvement over comparable systems.

a system executes a predefined task with limited contemporaneous human intervention. In augmentation, the system expands human information-processing capacity while a person retains the focal choice. In delegation, the system receives discretion to choose or act within a domain. These configurations can coexist in one process: an AI may automate data collection, augment diagnosis, select a priority class and return an exception to a human reviewer [1, 7].

The distinction matters because complementarity is not automatic. Humans retain strengths in contextual interpretation, moral judgment and reasoning under novelty, while AI can process large data volumes, apply rules consistently and detect statistical patterns. Collaboration improves performance only when the parties possess substantively different information or capabilities and the process allows those differences to be integrated [5, 27]. Productive delegation also depends on managers' metaknowledge of the relative limits of human and AI capabilities [35]. A broad

label such as AI-driven or human-in-the-loop conceals where discretion actually resides. Selective delegation therefore requires an account of the specific rights transferred at each stage.

Delegation preferences also reflect controllability, explainability and the moral character of the task. Decision-affected stakeholders may demand more human involvement than decision makers themselves [28]. Governance must therefore address both operational performance and the legitimacy of the authority arrangement.

3.2 AI reliability, calibrated reliance and meaningful control

AI capability and AI reliability are distinct. Capability is the task-specific ability to generate useful predictions, analyses or actions relative to human alternatives. Reliability is the extent to which that capability produces stable, calibrated and reproducible performance within a validated domain across time, relevant subgroups and foreseeable operating

conditions. A highly capable model may be unreliable after distribution shift, under sparse data or in an unvalidated subgroup. Reliability is therefore an antecedent and dynamic contingency of right allocation rather than a direct outcome of governance.

Human trust should be calibrated rather than maximized. Empirical research on human trust in AI shows that trust is shaped by perceived competence, benevolence and integrity, and that misplaced trust—whether excessive or deficient—degrades joint decision quality [36]. Explanations can support calibration when they reveal uncertainty or diagnostic reasons, but they can also increase persuasive force without improving understanding [30]. Cognitive forcing interventions reduce overreliance by requiring an initial judgment or comparison before AI advice is accepted [31]. Evidence from public-sector decision-making shows both automation bias and selective adherence to algorithmic advice [37]. Explanation design must also align with stakeholder roles and can be strengthened by social transparency about prior human use and organizational context [38, 39]. More broadly, human-AI collaboration requires identifying the structural conditions under which combining human and AI inputs improves decision outcomes, even when neither party alone holds a clear performance advantage [40]. Meaningful control consequently requires epistemic access, time to intervene, formal override authority and organizational protection for justified challenge [10, 11].

Governance contributes by producing evidence about reliability through validation, monitoring, incident analysis and feedback. It cannot convert an unreliable model into a reliable one by policy alone. When reliability is low or uncertain, the permissible scope of selection and execution rights should contract, while approval, escalation and monitoring should intensify. When reliability is stable, auditable and limited to a bounded domain, operative rights may expand without eliminating human veto and oversight.

3.3 Organizational governance and lifecycle responsibility

Responsible AI governance translates values into organizational arrangements. Reviews of organizational AI governance identify structural, procedural and relational practices [8, 19]. Best-practice and agenda-setting studies emphasize coordination, implementation barriers and knowledge gaps [20, 21]. More recent syntheses document

continuing problems of fragmented ownership, weak incentives and limited operational capability [22, 23]. Sociotechnical accounts locate ethical outcomes in the interaction among human agency, technological affordances and organizational structures [41]. Studies of ethics tools and corporate implementation further show that checklists and principles are weak when detached from ownership, incentives and escalation channels [42, 43]. Common barriers include unclear mandates, inadequate resources and competing objectives [44]. A lifecycle perspective is essential because problem formulation, data selection, model development, implementation, use, monitoring and retirement all shape the final decision [9]. A person who approves the last step may therefore lack control over important earlier choices.

Continuous auditing seeks evidence across the configured decision system rather than a one-time model review [24]. Ethics-based and emerging AI audit research likewise emphasizes scope, evidence and follow-up [25, 26]. Responsibility must extend across design, process, downstream outcomes, governance oversight and remediation. External providers may carry contractual and technical duties, but deploying enterprises cannot fully transfer their responsibility for governing material effects in the use contexts they select. Effective governance therefore requires that deploying enterprises maintain answerability across design, process and outcome domains—a standard that presupposes authority recognition, the capacity for interrogation, and meaningful limitation of power over deployed systems [32, 34].

4 An Integrated Framework of Decision Rights and Responsibility

4.1 Conceptual boundaries

The framework separates five constructs that are often blurred. Consequence severity is the magnitude of harm from an erroneous decision and is only one component of decision exposure. Decision exposure is the combined burden created by consequence severity, reversibility, ethical salience, uncertainty and stakeholder reach; it determines governance intensity rather than describing model performance. AI capability concerns what a system can do in a focal task, whereas AI reliability concerns whether that capability remains stable, calibrated and warranted within the validated domain over time. AI autonomy is the scope of operative decision rights actually delegated to AI and is therefore distinct from both

Table 3. Decision rights in enterprise human-AI collaborative decision-making.

Right	Core question	Illustrative AI role	Essential control
Information access	What data may be retrieved, combined or inferred?	Collects and synthesizes approved data	Access limits, data lineage and purpose controls
Recommendation	Who may generate or rank alternatives?	Ranks suppliers or proposes actions	Uncertainty display and alternative generation
Selection	Who chooses among approved options?	Chooses within a bounded set	Thresholds, domain validation and exception rules
Approval	Who activates a proposed choice before effect?	Normally withheld in exposed decisions	Time, competence and independent judgment
Veto	Who can stop, reject or suspend the process?	Triggers a stop condition	Protected override and restart authorization
Execution	Who implements the selected action?	Executes routine transactions within limits	Permission boundaries, logging and rollback
Escalation	When and to whom must a case be transferred?	Routes anomalies or contested cases	Named destination, response time and closure evidence

capability and technical sophistication.

4.2 Decision context, capability and role architecture

The framework begins with the decision, not the technology. As noted above, consequence severity is one component of decision exposure; together, five dimensions constitute the composite exposure burden. Consequence severity concerns the magnitude of potential harm. Reversibility concerns the feasibility and cost of correction. Ethical salience concerns rights, fairness, dignity, safety or public trust. Uncertainty concerns environmental stability and similarity to the validation domain. Stakeholder reach concerns whether consequences remain internal or extend to employees, customers, suppliers, communities or regulators. Exposure should also account for error propagation because small case-level errors may become material when automated at scale.

The enterprise then compares human and AI capability and assesses AI reliability. Relevant dimensions include data access and quality, predictive performance, contextual understanding, value judgment, adaptability and response speed. Reliability evidence should cover calibration, subgroup performance, robustness, drift, reproducibility and the completeness of monitoring. A capability advantage without reliable evidence does not justify broad operative authority.

Four role architectures follow. Human-led support uses AI to inform and recommend. Constrained augmentation allows AI to analyze, rank and prioritize while qualified humans retain exposed choices. Bounded automation permits AI selection

and execution within explicit thresholds and validated domains. Restricted delegation limits AI to analysis or simulation when proposed autonomy and decision exposure exceed the organization's assurance capacity. Role architecture may vary by process stage rather than apply to the workflow as a single label.

4.3 The decision-rights portfolio and protective safeguards

The paper distinguishes seven decision rights. Information-access rights authorize retrieval, combination or inference of data. Recommendation rights authorize proposing or ranking alternatives. Selection rights authorize choosing from an approved set. Approval rights authorize activating a proposed choice. Veto rights authorize stopping, rejecting or suspending the process. Execution rights authorize implementing the selected action. Escalation rights authorize or require transfer to another person or organizational level. These rights are separable and can be assigned to different actors, thresholds and model versions.

The portfolio is supplemented by protective safeguards rather than additional rights. Protected non-reliance operationalizes the human veto or override right when relevant evidence justifies departure from AI advice. Access to human reconsideration is implemented through escalation, contestability and remediation. Treating these safeguards as implementations of existing rights preserves the seven-right structure while making negative protection explicit. The seven decision rights, their core questions, illustrative AI roles, and essential controls are summarized in Table 3.

Table 4. Responsibility domains in enterprise human-AI decision governance.

Responsibility domain	Typical primary owner	Core duty	Required evidence
System design	Model owner and data owner	Define the problem, data, objectives, constraints and validation criteria	Model documentation, data lineage, validation and change records
Decision process	Business process owner	Ensure approved use, input quality, human review and escalation	Workflow configuration, training records and decision logs
Outcome stewardship	Business executive and risk function	Monitor downstream effects and act on harmful or unfair patterns	Outcome indicators, complaints, subgroup analysis and incident reports
Oversight	Executive committee or board-designated body	Set risk appetite, approve high-exposure uses and commission assurance	Risk acceptance, review minutes, assurance findings and remediation tracking
Remediation	Business owner with legal, HR or customer functions	Correct decisions, provide redress, compensate and learn	Appeal files, correction notices, compensation and control updates

4.4 The responsibility portfolio

Five responsibility domains are distinguished. System-design responsibility concerns problem definition, data, model objectives, constraints and validation. Decision-process responsibility concerns approved use, input quality, interpretation and contemporaneous review. Outcome-stewardship responsibility concerns downstream effects and harmful patterns. Oversight responsibility concerns risk appetite, role assignment, independent review and assurance. Remediation responsibility concerns correction, compensation, redress, learning and withdrawal where necessary.

Responsibility should be allocated according to effective control, competence and benefit. Effective control is the practical capacity to influence or prevent the relevant failure. Competence is the knowledge and resources required to exercise that control. Benefit is the organizational advantage gained from deployment. These criteria prevent responsibility dumping onto frontline users who lack authority and responsibility evasion by technical teams, vendors or executives who shape outcomes but remain outside review. The benefit criterion is particularly important for external providers: an enterprise that derives commercial advantage from deploying an AI system retains governance obligations that cannot be transferred to a vendor through contractual indemnification alone. Table 4 provides an overview of each responsibility domain, its typical primary owner, core duties, and the evidence required to discharge those duties.

4.5 The rights-control-responsibility alignment mechanism

The central explanatory mechanism is rights-control-responsibility alignment. First, allocating a decision right changes who possesses

discretion, evidence and the ability to intervene. Second, assigning responsibility creates an expectation that an actor will prevent, detect or remedy a class of failure. Third, governance quality depends on whether the responsible actor has the competence, authority and evidence required to discharge that expectation. Alignment exists when these elements coincide; misalignment exists when formal responsibility and effective control diverge.

Misalignment operates through three causal pathways. An epistemic deficit arises when the responsible actor cannot access model documentation, uncertainty or logs. An operational deficit arises when the actor can identify a problem but lacks veto, execution or escalation authority. An incentive deficit arises when an actor controls or benefits from the system while responsibility is placed elsewhere. These deficits make symbolic approval, delayed intervention and post-incident blame shifting more likely. The mechanism therefore explains governance failure even when the model is technically accurate and a human is formally present.

Alignment is dynamic. Reliability, data quality, organizational competence, scale and external expectations change over time. Evidence from validation, monitoring, incidents and appeals should trigger expansion or contraction of AI rights. Stable and auditable reliability may justify bounded automation for low-exposure decisions; drift, missing evidence or expanded use should automatically narrow operative rights and strengthen escalation. The overall architecture of the proposed enterprise human-AI decision governance framework, including its decision context, role architecture, decision-rights portfolio, responsibility domains and alignment mechanism, is illustrated in Figure 1.

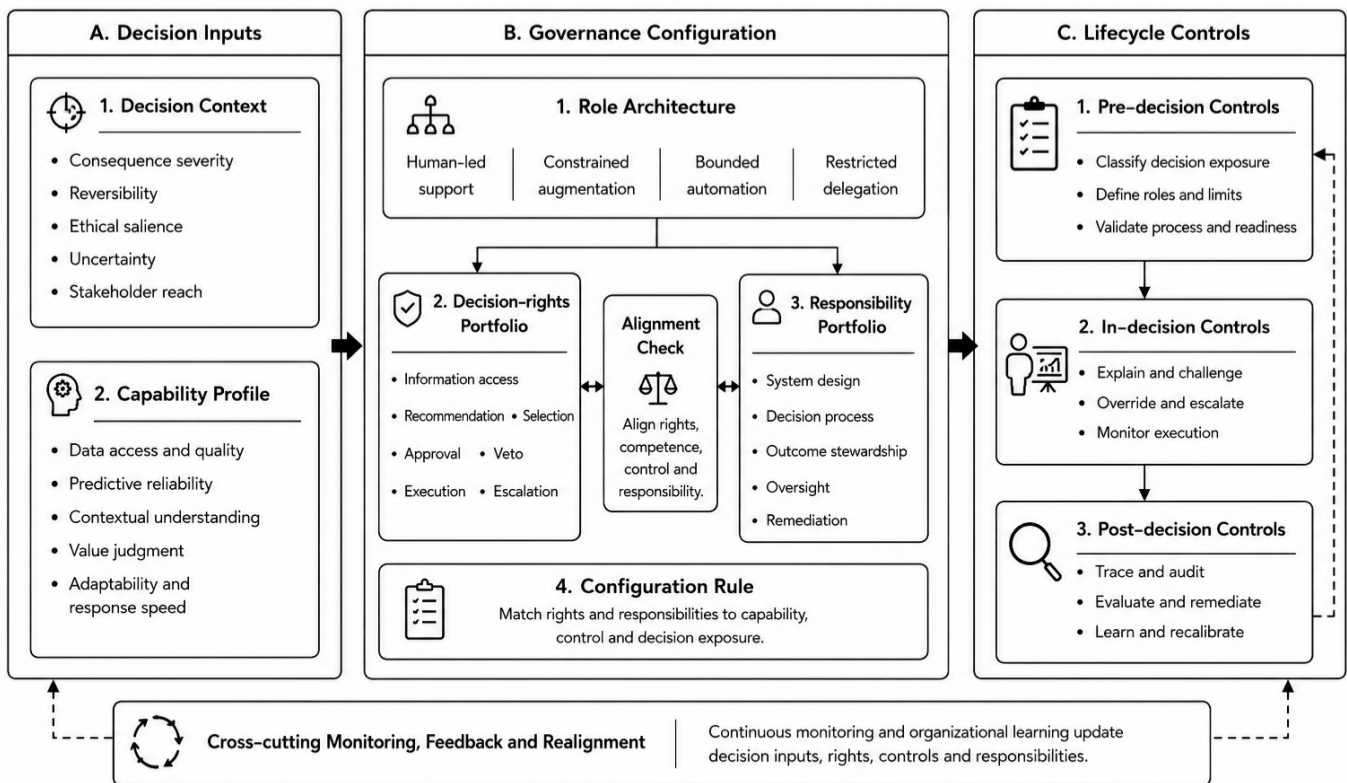


Figure 1. Enterprise human-AI decision governance architecture.

5 Normative Contingency Governance Design

Sections 5.1-5.5 state normative design principles. They describe the configuration required for governability, not a prediction that enterprises will necessarily adopt it. The principles translate the alignment mechanism into pre-decision, in-decision and post-decision controls.

5.1 Pre-decision governance

Pre-decision governance classifies exposure, defines roles and limits, and validates readiness. Each use case should be registered at the level of the focal decision with its exposure dimensions, reliability evidence, model version, stakeholders and named owners. A decision-rights map should identify who may access information, recommend, select, approve, veto, execute and escalate. Thresholds should be encoded in permissions and workflow rules rather than left as policy statements.

Validation should test the configured human-AI process, not the model in isolation. Evidence should include calibration, subgroup performance, edge cases, drift sensitivity, uncertainty communication and the ability of intended users to detect errors. Readiness requires decision-specific training, sufficient review time and protected challenge. A reviewer who faces

an excessive caseload does not provide meaningful control.

5.2 In-decision governance

In-decision governance converts information into action. Interfaces should communicate confidence, missing data, scope limitations and relevant counterevidence. Challenge mechanisms may include an independent initial judgment, mandatory comparison of alternatives, second review for disagreement and protected override. The system should record whether the human agreed, overrode, deferred or escalated without turning routine decisions into excessive administration.

Escalation should respond to low confidence, out-of-domain inputs, human-AI disagreement, stakeholder challenge or cumulative exposure limits. It must lead to a person with greater authority or expertise. Authorized business, risk and technical roles should be able to suspend automated execution, while restart authority and required evidence should be predefined.

5.3 Post-decision governance

Post-decision governance supports traceability, outcome evaluation, remediation and rights recalibration. Logs should connect input data,

model version, output, human action, override, execution and subsequent outcome. Monitoring should examine subgroup effects, error concentration, override patterns, distribution shift, complaints and downstream consequences rather than average predictive performance alone.

Audit should evaluate the configured decision system, including whether rights are enforced, reviewers are competent, exception cases are resolved appropriately, and providers supply required evidence. Redress should provide an accessible route to human reconsideration, correction and compensation. Incident and appeal evidence should update thresholds, training, model constraints and the underlying business objective.

5.4 Four governance archetypes and reliability conditions

Decision exposure and AI autonomy jointly yield four governance archetypes. AI reliability does not create a separate axis in the contingency matrix because it conditions whether a proposed position on the autonomy axis is permissible. Human-led support tolerates emerging reliability because AI outputs are non-binding. Bounded automation requires high, stable and auditable reliability within a validated domain. Constrained augmentation may use useful but insufficiently reliable AI because humans retain exposed choices. Restricted delegation applies when reliability is uncertain, unassurable or inadequate relative to exposure.

The four governance archetypes are distinguished by the level of decision exposure and the degree of AI autonomy. Low exposure permits higher AI autonomy (bounded automation), while high exposure restricts AI autonomy and requires stronger human retention of approval, veto and escalation rights (constrained augmentation). When both exposure and proposed autonomy are high, restricted delegation limits AI to analysis and simulation only.

The complete characteristics of each archetype, including AI and human roles, reliability conditions and core controls, are summarized in Table 5. The contingency matrix mapping these four archetypes by decision exposure and AI autonomy is illustrated in Figure 2.

5.5 Worked applications of the framework

The following applications show how the framework produces a more precise governance design than the

instruction to keep a human in the loop. The examples are illustrative rather than empirical cases. They apply every decision right, responsibility domain and evidence requirement to one low-exposure and one high-exposure decision.

5.5.1 Low-exposure application: routine inventory replenishment

Consider an AI system that replenishes standardized inventory for stable products. Consequence severity and ethical salience are low, transactions are reversible and stakeholder reach is limited. Reliability is high within validated product categories, but the system is not authorized for novel suppliers or exceptional demand shocks; bounded automation is therefore appropriate. The AI receives approved ERP, inventory and lead-time access, recommends order timing and quantity, selects within approved bounds and executes orders through restricted permissions. Human actors approve the policy and thresholds, retain veto authority and receive automatic escalations for low confidence, unusual demand, new suppliers or cumulative exposure. Model and data owners hold system-design responsibility; the procurement process owner governs approved use and exception handling; the supply-chain executive monitors stockouts, excess inventory and supplier effects; the AI governance or risk committee provides oversight; and the business owner reverses orders, corrects data and updates thresholds when remediation is required. Required evidence includes validation by product class, data lineage, transaction logs, exception rates, drift alerts and rollback records. This design avoids meaningless per-transaction clicks while preserving real veto, escalation and accountability.

5.5.2 High-exposure application: AI-assisted employee termination

Consider an AI system that identifies performance patterns relevant to an employee-termination process. Consequence severity, ethical salience and stakeholder reach are high, reversibility is limited and contextual uncertainty is substantial. Predictive evidence cannot capture all legal, interpersonal and fairness considerations, so reliability is insufficient for operative authority; restricted delegation, or constrained augmentation when AI only prioritizes evidence, is appropriate. The AI may access only approved and relevant HR data, summarize evidence and flag cases for review, but a qualified HR-business panel retains selection, approval, veto, execution and escalation. HR analytics and data owners define

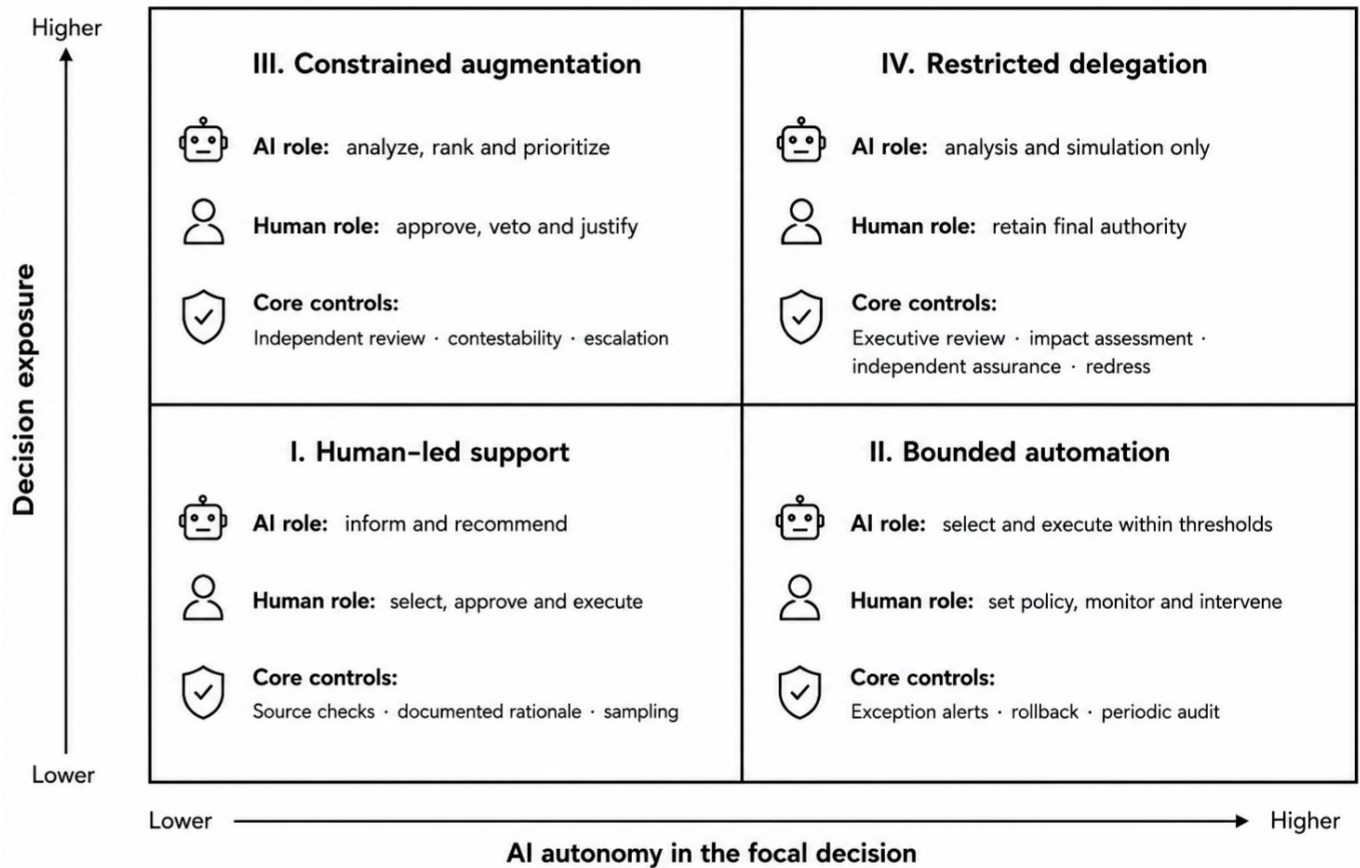


Figure 2. Contingency matrix of decision exposure and AI autonomy.

permissible purpose and validation boundaries; the HR process owner ensures independent review and non-retaliatory challenge; a business executive and HR risk function monitor fairness, complaints and workforce effects; an executive or board-designated body sets risk appetite and assurance requirements; and HR, legal and employee-relations functions provide appeal, correction, reinstatement or compensation where appropriate. Required evidence includes an impact assessment, subgroup testing, an individual reason record, independent review,

a durable audit trail and appeal and remediation records. Separating recommendation, selection, approval, veto, execution and appeal prevents the human reviewer from becoming a rubber stamp.

6 Empirically Testable Propositions

The propositions below describe relationships that can be observed and rejected; they do not state that enterprises ought to behave rationally. The dependent variables include the scope of AI rights, oversight quality, decision performance, recalibration

Table 5. Governance archetypes and reliability conditions.

Archetype	Exposure and AI autonomy	AI and human roles	Reliability condition	Core controls
Human-led support	Low exposure; low AI autonomy	AI informs and recommends; humans select, approve and execute	Emerging or uncertain reliability may be tolerated because outputs are non-binding	Source checks, documented rationale and sampling
Bounded automation	Low exposure; higher AI autonomy	AI selects and executes within thresholds; humans set policy, monitor and intervene	High, stable and auditable reliability within a validated domain	Exception alerts, rollback, cumulative limits and periodic audit
Constrained augmentation	High exposure; low AI autonomy	AI analyzes, ranks and prioritizes; humans approve, veto and justify	Useful performance may support advice but is insufficient for autonomous execution	Independent review, contestability and escalation
Restricted delegation	High exposure; high proposed autonomy	AI performs analysis and simulation only; humans retain final authority	Reliability is uncertain, unassessable or inadequate relative to exposure	Executive review, impact assessment, independent assurance and redress

speed, error discovery and perceived legitimacy. The proposition rows in Table 2 summarize the causal mechanisms, observable implications, moderators, competing explanations and rejection conditions.

Proposition 1. Greater uncertainty about AI reliability will be associated with narrower AI selection and execution rights and stronger human approval, escalation and monitoring, particularly when reversibility is low or error propagation is high.

Reliability uncertainty raises the expected cost of autonomous error. The relationship would not be supported if organizations systematically expand operative rights as reliability declines after controlling for exposure and economic pressure.

Proposition 2. Higher decision exposure will be associated with greater retention of human approval, veto and escalation rights and with higher governance intensity, even when AI predictive performance is superior.

Exposure increases legitimacy, accountability and remediation costs that predictive superiority does not eliminate. The effect should be stronger when ethical salience and stakeholder reach are high.

Proposition 3. Greater fit between allocated decision rights and task-specific human-AI capability asymmetries will be associated with higher joint decision quality and fewer preventable override or automation errors than allocations based on organizational status or technological enthusiasm.

Fit permits distinct information and competence to enter the stage where they are most useful. A rival explanation is that better-managed firms both allocate rights well and perform better; longitudinal or experimental designs can separate these effects.

Proposition 4. A larger misalignment between formal responsibility and effective control will be associated with more symbolic approval, slower intervention and greater post-incident responsibility shifting.

The mechanism is perceived inability to affect the outcome, reinforced by information and authority deficits. The relationship should be stronger under high reviewer workload and vendor opacity and would be rejected if severe mismatch coexists with consistently effective challenge and rapid intervention.

Proposition 5. Explanations will improve calibrated reliance only when they increase the responsible human's ability to discriminate between cases in which

AI advice should and should not be followed.

Diagnostic explanations improve discrimination; persuasive but non-diagnostic narratives may increase acceptance without improving correctness. The proposition is rejected when acceptance rises and discrimination does not, or when discrimination improves without explanation diagnosticity.

Proposition 6. Greater external-provider opacity will be associated with narrower permissible AI operative rights and stronger contractual audit, logging, change-notification and incident-cooperation requirements.

Opacity increases information asymmetry and weakens the enterprise's capacity to govern. The relationship should weaken when independent assurance, direct log access and switching options are available.

Proposition 7. More complete lifecycle traceability and continuous monitoring will be associated with faster and more accurate recalibration of AI rights after drift, incidents or use-context change.

Traceability identifies whether failure arose from data, model, workflow, human use or incentives. The effect depends on a governance actor having authority to suspend, modify or retire the system.

Proposition 8. More accessible contestability and effective redress will be associated with greater error discovery, faster correction, stronger perceived legitimacy and more organizational learning, especially when stakeholder reach is high.

Appeals provide external information that internal monitoring may miss. The proposition would be rejected if accessible, independent review produces no increase in detected errors, correction or legitimacy relative to otherwise comparable systems.

7 Discussion

7.1 Theoretical contributions

The first contribution is an explanatory mechanism rather than only a taxonomy. Rights-control-responsibility alignment explains why governance can fail despite accurate AI, formal human review and extensive policy. Delegating a right changes effective control and evidence access; responsibility becomes symbolic when the responsible actor lacks the means to intervene. Prior frameworks identify delegation, oversight or lifecycle controls

separately, whereas the present model connects them at the level of a focal decision.

The second contribution is analytical granularity. The seven-right portfolio distinguishes recommendation from selection, approval from veto and choice from execution. The five responsibility domains distinguish design from process, downstream stewardship, governance oversight and remediation. This decomposition reveals gaps that broad labels such as automation, augmentation or human-in-the-loop conceal.

The third contribution is the precise positioning of AI reliability as a governance contingency rather than a capability attribute or a direct outcome of governance. Reliability is neither identical to capability nor an automatic outcome of good governance. It is an antecedent and dynamic contingency that conditions permissible rights. Governance creates the evidence and feedback required to assess reliability and recalibrate authority. This distinction clarifies how technical performance enters organizational design without reducing governance to accuracy.

The fourth contribution is testability. The eight propositions identify dependent variables, mechanisms, moderators, competing explanations and falsification conditions. They permit experiments, surveys, field studies, process tracing and comparative cases to examine whether the proposed relationships hold rather than treating normative recommendations as predictions.

7.2 Managerial implications

Managers can use the framework as a design sequence. They should define the focal decision, assess exposure, compare capability, evaluate reliability, allocate the seven rights, assign the five responsibility domains and specify evidence before deployment. The worked cases show that meaningful control does not require a human click at every step. Low-exposure bounded automation can be governable when thresholds, veto, escalation and rollback are real. High-exposure decisions require substantive human selection, approval, contestability and redress.

Boards and executives should define risk appetite, prohibited uses, reliability standards and evidence requirements, escalation thresholds, and assurance expectations. Procurement contracts should secure audit access, logging, model-change notification, incident cooperation and exit options. Vendor opacity should narrow the scope of permissible AI

authority rather than justify attenuated governance requirements.

Human oversight should be evaluated as an operating capability through reviewer workload, disagreement resolution, override quality, escalation time, appeal outcomes and out-of-domain cases. Training without time, authority and psychological safety cannot produce meaningful control.

7.3 Boundary conditions, limitations and future research

The framework focuses on enterprise decisions in which an organization can define roles, permissions and evidence. It is less directly applicable to open consumer use of general-purpose AI or distributed ecosystems in which no organization controls the full interaction. Highly adaptive or opaque systems may require restricted delegation because reliability and alignment cannot be demonstrated.

The theory-building procedure is transparent but purposive rather than exhaustive. The corpus was coded by a single author, which constitutes a limitation: future research should replicate the category derivation with multiple independent coders, broader databases, and domain-specific literatures to assess inter-rater reliability and extend coverage. The worked applications are analytical illustrations rather than field evidence.

Future research can operationalize the seven rights as distinct permissions and test how they affect reliance, decision quality and legitimacy. Field studies can examine how responsibility maps shape incident response. Longitudinal research can trace how reliability evidence and incidents expand or contract AI rights. Cross-national studies can test whether regulation, professional norms and culture change acceptable delegation and contestability.

8 Conclusion

Enterprise human-AI decision-making is not governed adequately by the instruction to keep a human in the loop. Effective governance requires a precise account of what information AI may access; whether it may recommend, select, approve, veto, execute or escalate; who possesses effective control; and who is responsible for design, process, outcomes, oversight and remediation. This paper develops an auditable theory-building procedure and a rights-control-responsibility alignment mechanism. Decision exposure and AI autonomy

define four governance archetypes, while AI reliability conditions the permissible scope of operative rights. The framework becomes empirically useful through propositions that can be rejected and managerially useful through worked allocations of rights, responsibilities and evidence. Human-AI collaboration becomes governable when authority is explicit, reliability is evidenced, control is meaningful, and responsibility is anchored in actors who retain the capacity to shape outcomes.

Data Availability Statement

Data will be made available on request.

Funding

This work was supported without any funding.

Conflicts of Interest

The author declares no conflicts of interest.

AI Use Statement

The author declares that no generative AI was used in the preparation of this manuscript.

Ethical Approval and Consent to Participate

Not applicable.

References

- [1] Raisch, S., & Krakowski, S. (2021). Artificial intelligence and management: The automation-augmentation paradox. *Academy of Management Review*, 46(1), 192-210. [CrossRef]
- [2] Berente, N., Gu, B., Recker, J., & Santhanam, R. (2021). Managing artificial intelligence. *MIS Quarterly*, 45(3), 1433-1450. [CrossRef]
- [3] Feuerriegel, S., Hartmann, J., Janiesch, C., & Zschech, P. (2024). Generative AI. *Business & Information Systems Engineering*, 66(1), 111-126. [CrossRef]
- [4] Enholm, I. M., Papagiannidis, E., Mikalef, P., & Krogstie, J. (2022). Artificial intelligence and business value: A literature review. *Information systems frontiers*, 24(5), 1709-1734. [CrossRef]
- [5] Vaccaro, M., Almaatouq, A., & Malone, T. (2024). When combinations of humans and AI are useful: A systematic review and meta-analysis. *Nature Human Behaviour*, 8(12), 2293-2303. [CrossRef]
- [6] Puranam, P. (2021). Human-AI collaborative decision-making as an organization design problem. *Journal of Organization Design*, 10(2), 75-80. [CrossRef]
- [7] Baird, A., & Maruping, L. M. (2021). The next generation of research on IS use: A theoretical framework of delegation to and from agentic IS artifacts. *MIS Quarterly*, 45(1), 315-341. [CrossRef]
- [8] Papagiannidis, E., Mikalef, P., & Conboy, K. (2025). Responsible artificial intelligence governance: A review and research framework. *The Journal of Strategic Information Systems*, 34(2), 101885. [CrossRef]
- [9] Marabelli, M., Newell, S., & Handunge, V. (2021). The lifecycle of algorithmic decision-making systems: Organizational choices and ethical challenges. *The Journal of Strategic Information Systems*, 30(3), 101683. [CrossRef]
- [10] Green, B. (2022). The flaws of policies requiring human oversight of government algorithms. *Computer Law & Security Review*, 45, 105681. [CrossRef]
- [11] Te'eni, D., Yahav, I., & Schwartz, D. (2026). What it takes to control AI by design: human learning. *Ai & Society*, 41(1), 237-250. [CrossRef]
- [12] Tabassi, E. (2023). *Artificial Intelligence Risk Management Framework (AI RMF 1.0) (NIST AI 100-1)*. National Institute of Standards and Technology. [CrossRef]
- [13] International Organization for Standardization. (2023). *ISO/IEC 42001:2023 Information technology - Artificial intelligence - Management system*. ISO. <https://www.iso.org/standard/42001>
- [14] Jaakkola, E. (2020). Designing conceptual articles: four approaches. *AMS review*, 10(1), 18-26. [CrossRef]
- [15] Konsynski, B. R., Kathuria, A., & Karhade, P. P. (2024). Cognitive reapportionment and the art of letting go: A theoretical framework for the allocation of decision rights. *Journal of Management Information Systems*, 41(2), 328-340. [CrossRef]
- [16] Santoni de Sio, F., & Mecacci, G. (2021). Four responsibility gaps with artificial intelligence: Why they matter and how to address them. *Philosophy & Technology*, 34(4), 1057-1084. [CrossRef]
- [17] Stark, D., & Vanden Broeck, P. (2024). Principles of algorithmic management. *Organization Theory*, 5(2), 26317877241257213. [CrossRef]
- [18] Trunk, A., Birkel, H., & Hartmann, E. (2020). On the current state of combining human and artificial intelligence for strategic organizational decision making. *Business Research*, 13(3), 875-919. [CrossRef]
- [19] Mäntymäki, M., Minkkinen, M., Birkstedt, T., & Viljanen, M. (2022). Defining organizational AI governance. *AI and Ethics*, 2(4), 603-609. [CrossRef]
- [20] Papagiannidis, E., Enholm, I. M., Dremel, C., Mikalef, P., & Krogstie, J. (2023). Toward AI governance: Identifying best practices and potential barriers and outcomes. *Information Systems Frontiers*, 25(1), 123-141. [CrossRef]
- [21] Birkstedt, T., Minkkinen, M., Tandon, A., & Mäntymäki, M. (2023). AI governance: themes,

- knowledge gaps and future agendas. *Internet Research*, 33(7), 133-167. [CrossRef]
- [22] Batool, A., Zowghi, D., & Bano, M. (2025). AI governance: a systematic literature review. *AI and Ethics*, 5(3), 3265-3279. [CrossRef]
- [23] Sadek, M., Kallina, E., Bohné, T., Mougenot, C., Calvo, R. A., & Cave, S. (2025). Challenges of responsible AI in practice: scoping review and recommended actions. *AI & society*, 40(1), 199-215. [CrossRef]
- [24] Minkkinen, M., Laine, J., & Mäntymäki, M. (2022). Continuous auditing of artificial intelligence: a conceptualization and assessment of tools and frameworks. *Digital Society*, 1(3), 21. [CrossRef]
- [25] Mökander, J., Morley, J., Taddeo, M., & Floridi, L. (2021). Ethics-based auditing of automated decision-making systems: Nature, scope, and limitations. *Science and Engineering Ethics*, 27(4), 44. [CrossRef]
- [26] Schiff, D. S., Kelley, S., & Camacho Ibáñez, J. (2024). The emergence of artificial intelligence ethics auditing. *Big Data & Society*, 11(4), 20539517241299732. [CrossRef]
- [27] Hemmer, P., Schemmer, M., Kühl, N., Vössing, M., & Satzger, G. (2025). Complementarity in human-AI collaboration: Concept, sources, and evidence. *European Journal of Information Systems*, 34(6), 979-1002. [CrossRef]
- [28] Freisinger, E., & Schneider, S. (2025). Decoding decision delegation to artificial intelligence: A mixed-methods study on the preferences of decision-makers and decision-affected in surrogate decision contexts. *European Management Journal*, 43(6), 958-969. [CrossRef]
- [29] European Parliament and Council of the European Union. (2024). Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence and amending Regulations (Artificial Intelligence Act). *Official Journal of the European Union*, L 2024/1689. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>
- [30] Schemmer, M., Kuehl, N., Benz, C., Bartos, A., & Satzger, G. (2023, March). Appropriate reliance on AI advice: Conceptualization and the effect of explanations. In *Proceedings of the 28th International Conference on Intelligent User Interfaces* (pp. 410-422). [CrossRef]
- [31] Buçinca, Z., Malaya, M. B., & Gajos, K. Z. (2021). To trust or to think: cognitive forcing functions can reduce overreliance on AI in AI-assisted decision-making. *Proceedings of the ACM on Human-computer Interaction*, 5(CSCW1), 1-21. [CrossRef]
- [32] Novelli, C., Taddeo, M., & Floridi, L. (2024). Accountability in artificial intelligence: what it is and how it works. *AI & Society*, 39(4), 1871-1882. [CrossRef]
- [33] Minkkinen, M., Zimmer, M. P., & Mäntymäki, M. (2023). Co-shaping an ecosystem for responsible AI: Five types of expectation work in response to a technological frame. *Information Systems Frontiers*, 25(1), 103-121. [CrossRef]
- [34] Laux, J., Wachter, S., & Mittelstadt, B. (2024). Trustworthy artificial intelligence and the European Union AI act: On the conflation of trustworthiness and acceptability of risk. *Regulation & Governance*, 18(1), 3-32. [CrossRef]
- [35] Fügener, A., Grahl, J., Gupta, A., & Ketter, W. (2022). Cognitive challenges in human-artificial intelligence collaboration: Investigating the path toward productive delegation. *Information systems research*, 33(2), 678-696. [CrossRef]
- [36] Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of management annals*, 14(2), 627-660. [CrossRef]
- [37] Alon-Barkat, S., & Busuioc, M. (2023). Human-AI interactions in public sector decision making: "automation bias" and "selective adherence" to algorithmic advice. *Journal of Public Administration Research and Theory*, 33(1), 153-169. [CrossRef]
- [38] Langer, M., Oster, D., Speith, T., Hermanns, H., Kästner, L., Schmidt, E., ... & Baum, K. (2021). What do we want from Explainable Artificial Intelligence (XAI)?—A stakeholder perspective on XAI and a conceptual model guiding interdisciplinary XAI research. *Artificial intelligence*, 296, 103473. [CrossRef]
- [39] Ehsan, U., Liao, Q. V., Muller, M., Riedl, M. O., & Weisz, J. D. (2021). Expanding explainability: Towards social transparency in AI systems. In *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems* (pp. 1-19). ACM. [CrossRef]
- [40] Choudhary, V., Marchetti, A., Shrestha, Y. R., & Puranam, P. (2025). Human-AI ensembles: When can they work? *Journal of Management*, 51(2), 536-569. [CrossRef]
- [41] Heyder, T., Passlack, N., & Posegga, O. (2023). Ethical management of human-AI interaction: Theory development review. *The Journal of Strategic Information Systems*, 32(3), 101772. [CrossRef]
- [42] Ayling, J., & Chapman, A. (2022). Putting AI ethics to work: are the tools fit for purpose?. *AI and Ethics*, 2(3), 405-429. [CrossRef]
- [43] Ibáñez, J. C., & Olmeda, M. V. (2022). Operationalising AI ethics: how are companies bridging the gap between practice and principles? An exploratory study. *AI & Society*, 37(4), 1663-1687. [CrossRef]
- [44] Morley, J., Kinsey, L., Elhalal, A., Garcia, F., Ziosi, M., & Floridi, L. (2023). Operationalising AI ethics: Barriers, enablers and next steps. *AI & Society*, 38(1), 411-423. [CrossRef]