

Journal of Social Science and Human Research Studies

Volume: 02 Issue: 05 May-2026

Decision Legitimacy in AI-Enabled Organizations: A Multilevel Framework of Algorithmic Authority and Human Judgment

Ashi TENDERİS

Independent Researcher, İstanbul, Turkey. ORCID ID: 0000-0001-6312-7832

Published Date: May 21, 2026

Article DOI: 10.65150/EP-jsshrs/V2E5/2026-14

ABSTRACT: This study develops a theoretical framework explaining decision legitimacy in AI-enabled organizations as an emergent outcome of the interplay between algorithmic authority and human judgment. Drawing on institutional theory, socio-technical systems, and behavioral decision-making, it identifies procedural, distributive, and cognitive legitimacy as key dimensions. While AI enhances efficiency and consistency, legitimacy depends on transparency, explainability, and perceived fairness. Algorithmic authority may both strengthen and undermine legitimacy, whereas human judgment remains essential for contextual interpretation and accountability. The study advances multilevel propositions and outlines a research agenda for examining legitimacy in hybrid human–AI decision systems.

KEYWORDS: Decision Legitimacy; Algorithmic Authority; Human Judgment; AI-Enabled Organizations; Theoretical Framework

INTRODUCTION

The pervasive integration of artificial intelligence within organizational structures has reconfigured traditional decision-making paradigms, extending AI from a supportive tool to an active participant in decision formulation, implementation, and monitoring (Shi, 2026). This transformation necessitates a critical rethinking of established organizational theories, particularly due to the fundamental differences between human cognition and algorithmic processing (Shrestha et al., 2019). As AI increasingly participates in managerial tasks, important questions emerge regarding the allocation of decision rights within hybrid human–AI systems (Westover, 2025). In this context, traditional approaches grounded in human intuition and experience are increasingly challenged by AI's data-driven reasoning capabilities, which can simultaneously enhance and potentially displace human judgment (Upase, 2026).

Within this evolving landscape, organizations face a fundamental challenge in maintaining decision legitimacy as AI systems—often described as “invisible managers” or “shadow executives”—increasingly shape strategic and operational outcomes without clear accountability structures (Gharaibeh, 2026; Badgami, 2025). This development creates a significant “responsibility gap,” requiring new governance models suited to the emerging “algocratic” era (“Trust and Responsibility in AI-Driven Leadership: Accountability for AI Decisions,” 2026). Consequently, the increasing delegation of decision authority to AI raises pressing concerns regarding organizational legitimacy and alignment with broader societal values (Alhosani & Alhashmi, 2024).

In AI-enabled organizations, decision legitimacy depends on balancing the efficiency and predictive power of AI with human values, ethical considerations, and accountability requirements (Demirsoy, 2025; Gharaibeh, 2026). This necessitates a deeper understanding of how human actors perceive, interpret, and accept algorithmic influence in managerial processes, particularly as AI systems begin to perform leadership-like functions (Gladden et al., 2022; Haesevoets et al., 2021). Accordingly, human interpretive judgment remains central to organizational coordination and governance, especially in contexts where algorithmic outputs require contextualization and sensemaking (Hillebrand et al., 2025; Raisch & Krakowski, 2021).

Building on these gaps, this study develops a comprehensive theoretical framework of decision legitimacy in AI-enabled organizations, conceptualized as an emergent outcome of the dynamic interplay between algorithmic authority and human judgment. Rather than treating legitimacy as a static property of decisions or systems, the proposed framework positions it as a socially constructed and multilevel phenomenon shaped by technological, organizational, and cognitive dynamics. Drawing on institutional theory, socio-technical systems thinking, and behavioral decision-making literature, the framework identifies three core dimensions of legitimacy in AI contexts: procedural legitimacy (perceived fairness and transparency of decision processes), distributive legitimacy (perceived fairness of outcomes), and cognitive legitimacy (the extent to which AI-driven decisions are understood, accepted, and taken for granted). These dimensions are continuously reshaped as algorithmic systems become more deeply embedded in decision-making processes traditionally reserved for human actors.

Therefore, a multilevel framework is required to explain how algorithmic authority and human judgment interact to shape decision legitimacy across individual, group, and organizational levels (Esch, 2026). This requires integrating management theory, information systems research, and organizational ethics to conceptualize responsible AI as a dynamic capability rather than a static technological artifact (Akbarighatar, 2024; Hossain et al., 2025a, 2025b). Within this interaction, algorithmic authority does not exert a uniform influence on legitimacy perceptions; rather,

License: This is an open access article under the CC BY 4.0 license: <https://creativecommons.org/licenses/by/4.0/>

its effects are contingent on conditions such as transparency, explainability, perceived fairness, and accountability structures. At the same time, human judgment plays a critical role in providing contextual interpretation, ethical evaluation, and responsibility attribution, thereby shaping whether AI-supported decisions are perceived as legitimate or contested.

Accordingly, key trade-offs—such as efficiency versus fairness and transparency versus control must be continuously managed through explainability practices, human–AI collaboration, and accountability mechanisms (Reddy et al., 2026). Multilevel governance structures are therefore essential to address ethical risks such as bias, opacity, and diffused accountability while preserving meaningful human oversight in strategic decision-making (Choung et al., 2023; Mohammed et al., 2025; Pöhler et al., 2024). By incorporating transparency mechanisms, fairness audits, and structured human–algorithm complementarity, organizations can sustain normative legitimacy while leveraging the analytical power of AI systems (Batool et al., 2025; Oladele et al., 2024). Ultimately, such governance architectures position responsible AI not merely as a technical requirement but as a core organizational capability that enhances both resilience and legitimacy in hybrid decision systems (Madanchian & Taherdoost, 2025).

Based on this conceptualization, the study develops a set of theoretical propositions outlining the conditions under which hybrid human–AI decision-making structures enhance or undermine legitimacy across individual, organizational, and societal levels. These propositions aim to guide future empirical research and advance theory development in this emerging domain. Finally, the paper is structured as follows: the next section reviews the relevant literature on AI in organizational decision-making and legitimacy, followed by the presentation of the proposed theoretical framework and propositions. The paper concludes by outlining a future research agenda and discussing implications for the design of responsible, transparent, and legitimate AI-enabled decision systems.

2. LITERATURE REVIEW

2.1. Organizational Legitimacy and Decision Legitimacy

Organizational legitimacy is a foundational construct in institutional theory, defined as the generalized perception that organizational actions are desirable, proper, and appropriate within socially constructed systems of norms and beliefs (Suchman, 1995). Early institutional scholarship conceptualized legitimacy as a macro-level survival mechanism shaped by institutional isomorphism and conformity pressures (DiMaggio & Powell, 1983; Meyer & Rowan, 1977). However, subsequent research has increasingly emphasized the micro-foundations of legitimacy, arguing that legitimacy is not only structurally conferred but also cognitively constructed through individual-level evaluations (Bitektine & Haack, 2015; Tost, 2011).

Within this evolving literature, decision legitimacy has emerged as a distinct construct focusing on the perceived appropriateness of specific organizational decisions rather than organizational entities as a whole. Decision legitimacy is typically grounded in procedural justice, distributive justice, and interactional fairness (Colquitt et al., 2001; Greenberg, 1990). Empirical evidence consistently shows that individuals may accept unfavorable outcomes if the decision-making process is perceived as fair and transparent (Tyler, 2006; Tyler & Blader, 2003). This suggests that legitimacy is not outcome-dependent but process-sensitive, relying heavily on perceived fairness and trust in decision authority. Importantly, legitimacy is not static but socially constructed and context-dependent. It evolves through institutionalization processes in which repeated actions become taken-for-granted norms (Suddaby et al., 2017; Scott, 2014). This dynamic nature becomes particularly complex when decision authority shifts from human actors to algorithmic systems.

2.2. Algorithmic Decision-Making in Organizations

The increasing integration of artificial intelligence into organizational decision-making has fundamentally altered traditional governance structures. Algorithmic decision-making (ADM) refers to the use of computational systems that process large-scale data to generate or support decisions traditionally made by humans (Kellogg et al., 2020; Raisch & Krakowski, 2021). ADM systems are widely adopted in areas such as recruitment, performance evaluation, credit scoring, and resource allocation due to their efficiency, scalability, and perceived objectivity.

However, while ADM is often associated with rationality and neutrality, critical scholarship challenges this assumption by demonstrating that algorithms embed structural biases derived from historical data and design choices (Barocas & Selbst, 2016; Noble, 2018). Consequently, algorithmic outputs are not neutral but reflect socio-technical assemblages shaped by human design and institutional context (Mittelstadt et al., 2016; Burrell, 2016).

Furthermore, ADM systems disrupt conventional accountability structures. Unlike human decision-makers, algorithms lack intentionality and moral agency, creating a “responsibility gap” in organizational decision processes (Danaher et al., 2017; Kroll et al., 2017). This diffusion of responsibility complicates legitimacy attribution, as stakeholders struggle to determine whether accountability lies with developers, organizations, or the algorithm itself.

2.3. Algorithmic Authority versus Human Judgment

A central tension in the literature concerns the comparative authority of algorithmic systems and human decision-makers. Human judgment has traditionally been valued for its contextual sensitivity, ethical reasoning, and adaptive flexibility (Highhouse, 2008; Kahneman, 2011). In contrast, algorithms are valued for consistency, scalability, and resistance to cognitive bias (Logg et al., 2019).

This contrast has led to the conceptualization of algorithmic authority, defined as the degree to which decision legitimacy is attributed to algorithmic outputs rather than human judgment (Kellogg et al., 2020). Empirical studies suggest that algorithmic authority is conditional rather than absolute. Individuals tend to prefer algorithmic decisions in structured, data-rich environments, whereas human judgment is preferred in ambiguous or morally sensitive contexts (Lee, 2018; Castelo et al., 2019).

From a cognitive perspective, this preference structure is shaped by heuristic processing, where individuals associate algorithms with objectivity and humans with subjectivity (Kahneman, 2011). From an institutional perspective, legitimacy is embedded in normative expectations about appropriate authority distribution (Tost, 2011; Suddaby et al., 2017). The interplay between these mechanisms highlights the need for a multilevel framework that integrates cognitive and institutional dimensions of algorithmic authority.

2.4. Transparency, Fairness, and Trust in Algorithmic Systems

Transparency is widely considered a key determinant of algorithmic legitimacy, defined as the extent to which algorithmic logic and decision rules are understandable to stakeholders (Ananny & Crawford, 2018; Burrell, 2016). However, transparency is inherently constrained by the complexity of machine learning systems, often resulting in “black box” decision processes that limit interpretability (Zerilli et al., 2019). Fairness is another central dimension of algorithmic legitimacy, encompassing both procedural and distributive justice considerations. Research demonstrates that perceived fairness significantly influences acceptance of algorithmic decisions, particularly in high-stakes domains such as employment and justice systems (Binns et al., 2018; Lee, 2018). Importantly, fairness perceptions are not solely determined by outcomes but also by the perceived legitimacy of the decision process.

Trust functions as a mediating mechanism between transparency, fairness, and legitimacy. Trust in algorithmic systems is typically grounded in perceived competence but is highly sensitive to perceived errors and lack of explainability (Mayer et al., 1995; Glikson & Woolley, 2020). Notably, research on algorithm aversion shows that individuals may reject algorithmic systems after observing even minor errors, despite their superior average performance (Dietvorst et al., 2015). This indicates that trust in algorithmic systems is fragile and context-dependent.

2.5. The Paradox of Algorithmic Legitimacy

A key insight emerging from the literature is the paradoxical nature of algorithmic legitimacy. On one hand, algorithmic systems enhance legitimacy through perceived objectivity, consistency, and data-driven decision-making (Logg et al., 2019). On the other hand, they undermine legitimacy due to opacity, lack of accountability, and perceived dehumanization (Glikson & Woolley, 2020; Bitektine & Haack, 2015). This duality generates what can be conceptualized as an algorithmic legitimacy paradox, in which the same system simultaneously strengthens and weakens legitimacy depending on contextual interpretation. Existing research has largely treated these effects separately, without fully integrating their interactive and dynamic nature.

The paradox is further intensified by the socio-technical nature of algorithmic systems, which are neither purely technical nor purely social but hybrid constructs shaped by institutional, organizational, and cognitive forces (Mittelstadt et al., 2016; Faraj et al., 2018). This hybridity complicates traditional legitimacy frameworks and necessitates a more integrative theoretical approach.

2.6. Toward a Multilevel Theory of Decision Legitimacy

Existing literature remains fragmented across analytical levels. Micro-level studies focus on individual cognition and perception (Tost, 2011), meso-level studies examine organizational implementation practices (Kellogg et al., 2020), while macro-level research addresses institutional norms and legitimacy structures (DiMaggio & Powell, 1983; Scott, 2014). However, these levels are often treated independently, limiting theoretical integration.

A multilevel perspective is essential for understanding decision legitimacy in AI-augmented organizations. At the micro level, individuals interpret algorithmic decisions through cognitive heuristics and fairness perceptions. At the meso level, organizational governance structures mediate how algorithms are deployed, explained, and legitimized. At the macro level, institutional environments define normative expectations regarding acceptable forms of decision authority (Deephhouse et al., 2017; Suddaby et al., 2017). Integrating these levels allows for a more comprehensive theory of decision legitimacy that captures the interaction between algorithmic authority and human judgment. Such a framework advances current literature by moving beyond static, single-level explanations toward a dynamic, multilevel understanding of legitimacy construction in AI-augmented organizations.

3. CONCEPTUAL FRAMEWORK: DECISION LEGITIMACY IN AI-ENABLED ORGANIZATIONS

3.1. Algorithmic Authority

The growing integration of artificial intelligence into organizational decision-making processes has given rise to a new form of authority—algorithmic authority. Unlike traditional forms of authority grounded in hierarchy or professional expertise, algorithmic authority derives its legitimacy from perceived objectivity, data-driven rationality, and computational accuracy (Gillespie, 2014; Beer, 2017). AI systems are often viewed as impartial and consistent decision-makers, which can enhance trust in decision outcomes (Brynjolfsson & McAfee, 2017). However, this perception is not unproblematic.

Algorithmic authority is inherently shaped by underlying data structures, model assumptions, and design choices, which may remain opaque to users (Pasquale, 2015). This opacity introduces challenges related to transparency and explainability, potentially undermining the perceived legitimacy of decisions (Burrell, 2016). Consequently, algorithmic authority operates as a double-edged construct: it can both reinforce and erode decision legitimacy depending on how it is perceived and enacted within organizational contexts.

3.2. Human Judgment

Despite the increasing reliance on AI systems, human judgment remains a central component of organizational decision-making. Human actors contribute contextual awareness, ethical reasoning, and interpretive flexibility—capabilities that are not easily replicated by algorithmic systems (Kahneman, 2011; Simon, 1997). In AI-enabled environments, human judgment plays a critical role in validating, contesting, or overriding algorithmic outputs (Davenport & Kirby, 2016).

Moreover, human involvement is essential for assigning responsibility and accountability, particularly in complex or morally ambiguous decisions (Binns, 2018). The presence (or absence) of human oversight significantly influences how stakeholders perceive the fairness and legitimacy of decisions. Therefore, rather than being replaced, human judgment interacts dynamically with algorithmic authority, shaping the overall legitimacy of decision processes.

3.3. Dimensions of Decision Legitimacy

Building on prior literature, this study conceptualizes decision legitimacy as a multidimensional construct consisting of procedural, distributive, and cognitive dimensions (Suchman, 1995).

Procedural legitimacy refers to the perceived fairness, transparency, and consistency of decision-making processes. In AI-enabled contexts, procedural legitimacy is closely tied to factors such as explainability and the visibility of decision rules (Burrell, 2016).

Distributive legitimacy concerns the perceived fairness of decision outcomes, including whether results are equitable and unbiased. Algorithmic bias and data-driven inequalities pose significant challenges to distributive legitimacy (O'Neil, 2016).

Cognitive legitimacy reflects the extent to which decisions are understandable, acceptable, and taken for granted by organizational members (Suchman, 1995). The complexity and opacity of AI systems can hinder the development of cognitive legitimacy, particularly when users lack sufficient understanding of how decisions are produced (Pasquale, 2015).

These dimensions are interrelated and jointly shape the overall perception of legitimacy in AI-supported decision-making environments.

3.4. An Integrative Framework

Integrating the above elements, this study proposes that decision legitimacy in AI-enabled organizations emerges from the dynamic interplay between algorithmic authority and human judgment across multiple levels of analysis.

Algorithmic authority influences legitimacy by enhancing efficiency and consistency, but may simultaneously undermine it through opacity and perceived lack of accountability (Pasquale, 2015; Burrell, 2016). Human judgment, in turn, acts as a moderating and interpretive mechanism that can either reinforce or challenge algorithmic outputs (Davenport & Kirby, 2016). The balance between these two forces determines how legitimacy is constructed and sustained. Furthermore, this relationship is contingent upon key boundary conditions, including transparency, explainability, perceived fairness, and accountability structures (Binns, 2018). These conditions shape whether algorithmic authority is accepted as a legitimate basis for decision-making or resisted as a source of uncertainty and distrust.

Accordingly, decision legitimacy should be understood not as a fixed outcome, but as an evolving and context-dependent phenomenon that is continuously negotiated within hybrid human-AI decision systems.



Figure 1. A Multilevel Framework of Decision Legitimacy in AI-Enabled Organizations

4. THEORETICAL PROPOSITIONS

Building on the proposed conceptual framework, this study develops a set of theoretical propositions that explain how the interaction between algorithmic authority and human judgment shapes decision legitimacy in AI-enabled organizations.

P1. Algorithmic authority is positively associated with perceived procedural legitimacy when decision processes are transparent and explainable.

P2. Algorithmic authority is negatively associated with procedural legitimacy when decision processes are perceived as opaque or incomprehensible.

P3. Algorithmic authority is positively associated with distributive legitimacy when outcomes are perceived as unbiased and data-driven.

P4. Algorithmic authority is negatively associated with distributive legitimacy when algorithmic bias or data inequality is perceived.

P5. Human judgment positively influences cognitive legitimacy by enhancing the interpretability and contextual understanding of AI-supported decisions.

P6. Human involvement in decision-making processes strengthens perceptions of accountability, thereby increasing overall decision legitimacy.

P7. The positive effect of algorithmic authority on decision legitimacy is strengthened when complemented by human judgment.

P8. The absence of human oversight weakens the relationship between algorithmic authority and perceived legitimacy.

P9. The relationship between algorithmic authority and decision legitimacy varies across levels (individual, organizational, societal).

P10. Hybrid human-AI decision systems are more likely to generate sustained legitimacy when transparency, fairness, and accountability conditions are jointly satisfied.

5. DISCUSSION

This study develops a multilevel theoretical framework to explain how decision legitimacy emerges in AI-enabled organizational contexts, where algorithmic authority and human judgment coexist, interact, and occasionally conflict. Building on the proposed theoretical propositions, the findings suggest that legitimacy in such environments is neither inherently derived from technological sophistication nor solely from human

oversight, but rather from the dynamic alignment between algorithmic outputs and socially constructed expectations of fairness, accountability, and transparency. This interpretation is consistent with foundational perspectives that conceptualize legitimacy as a generalized perception of appropriateness within socially constructed systems of norms and values (Suchman, 1995).

A central insight of this study is that algorithmic authority does not operate in isolation; instead, it is continuously interpreted, contested, and validated by human actors across organizational levels. While algorithms provide consistency, efficiency, and data-driven rationality, these features alone are insufficient to ensure legitimacy. Rather, legitimacy emerges when organizational members perceive algorithmic decisions as understandable, justifiable, and aligned with shared norms and values (Suchman, 1995; Meyer & Rowan, 1977). This aligns with emerging research emphasizing that procedural fairness, transparency, and explainability are essential for fostering trust in AI systems (Dwork et al., 2012; Wachter, Mittelstadt, & Russell, 2017).

Furthermore, the framework demonstrates that human judgment retains a critical yet transformed role in AI-enabled decision-making processes. Rather than being displaced, human actors function as sensemakers, moderators, and legitimizers of algorithmic outputs. This reflects a shift toward hybrid governance structures in which legitimacy is co-produced through ongoing interaction between human and algorithmic agents. Prior research similarly highlights the importance of human-centered design in ensuring that AI systems remain interpretable, accountable, and aligned with organizational objectives (Shneiderman, 2020; Davenport & Ronanki, 2018). However, tensions may arise when algorithmic recommendations conflict with human intuition or ethical considerations, potentially leading to legitimacy breakdowns or resistance behaviors (Yeung, 2018; Mittelstadt et al., 2016).

At the multilevel perspective, legitimacy is shaped through the interaction of individual cognition, organizational structures, and institutional environments. At the individual level, cognitive trust and perceived procedural fairness influence acceptance of AI outputs (Dwork et al., 2012). At the organizational level, governance mechanisms, transparency practices, and accountability structures determine how algorithmic decisions are communicated and justified (Davenport & Ronanki, 2018). At the institutional level, regulatory frameworks, cultural expectations, and societal norms define the boundaries of acceptable algorithmic authority (Meyer & Rowan, 1977; Yeung, 2018). Together, these levels underscore that legitimacy is an emergent, dynamic, and context-dependent process rather than a static organizational attribute.

From a future research perspective, this study highlights the importance of examining subjective experiences and cognitive construals of individuals interacting with AI systems to better understand trust and accountability in evolving organizational contexts (Cabiddu et al., 2022; Ossa et al., 2024; Schmidt et al., 2025). To enhance generalizability, future studies should employ large-scale quantitative research designs and extend analyses beyond European contexts to capture cultural and institutional variation in AI adoption (Stahl et al., 2021). Comparative sectoral studies—particularly in FinTech and supply chain management—are also necessary to examine how industry-specific conditions shape AI bias and ethical risks (Sharabati et al., 2024).

Longitudinal research is essential to capture how AI reshapes organizational processes, structures, and performance over time, as well as how human–AI collaboration evolves across stages of technological integration (Leoni et al., 2024). Future studies should also investigate AI's role in complex decision domains such as negotiation and conflict resolution, where its capabilities remain limited, to better understand hybrid intelligence development (Sudeeptha et al., 2024). In parallel, deeper attention must be given to ethical concerns such as algorithmic bias, employee marginalization, and decision opacity, reinforcing the need for robust frameworks of fairness, accountability, and transparency in AI systems (Wójcik, 2025).

Moreover, understanding organizational culture for effective human–AI teaming is critical, particularly cultural conditions that foster trust and adaptive collaboration (Dima et al., 2024). Employee attitudes toward AI—including ambivalence and resistance—should also be examined as key determinants of implementation success (Chiu et al., 2021). Future research should further assess the long-term effects of externally sourced AI systems, distinguishing between customizable and off-the-shelf solutions and comparing outcomes across developing and user firms (Malik et al., 2020). This includes analyzing impacts on resource allocation, organizational restructuring, and human resource practices (Enholm et al., 2021; Wójcik, 2024).

Finally, longitudinal studies on human–AI relationships are needed to understand how these interactions evolve over time, including their developmental stages and ethical implications (Koponen et al., 2023). Particular attention should be given to the long-term consequences of repeated interaction with AI supervisors, especially regarding compliance with potentially unethical instructions (Lanz et al., 2023). Overall, decision legitimacy in AI-enabled organizations depends not only on algorithmic performance but on the extent to which AI systems are embedded within transparent, accountable, and human-centered governance structures; otherwise, organizations risk diminished trust, resistance, and weakened effectiveness of AI-driven transformation processes (Davenport & Ronanki, 2018; Shneiderman, 2020).

6. THEORETICAL IMPLICATIONS

This study offers several important theoretical contributions to the emerging literature on artificial intelligence (AI) in organizations, decision-making, and organizational legitimacy. By developing a multilevel framework that integrates algorithmic authority and human judgment, the study advances existing theory in three key ways: by extending legitimacy theory into AI-enabled contexts, by reconceptualizing authority as hybrid and relational, and by introducing a multilevel perspective that bridges micro, meso, and macro domains.

First, the study contributes to organizational legitimacy theory by extending its scope to algorithmically mediated environments. Traditional conceptualizations of legitimacy emphasize social acceptance, institutional conformity, and normative alignment within human-centered systems (Suchman, 1995; Meyer & Rowan, 1977). However, these perspectives largely assume that authority is exercised by human actors or formal institutions. This study challenges that assumption by demonstrating that algorithms themselves can function as quasi-authoritative agents, shaping perceptions of appropriateness and influencing decision outcomes. In doing so, it introduces the concept of *algorithmic legitimacy* as a distinct yet interdependent dimension of organizational legitimacy. This extension responds to recent calls for rethinking core organizational constructs in light of digital transformation and algorithmic governance (Mittelstadt et al., 2016; Yeung, 2018).

Second, the study advances theory by reconceptualizing authority as a hybrid and relational construct. Rather than viewing algorithmic authority and human judgment as competing or mutually exclusive, the proposed framework positions them as co-evolving and interdependent sources of influence. This perspective challenges deterministic narratives that portray AI as either replacing or subordinating human decision-making. Instead, it aligns with emerging research emphasizing human–AI complementarity and augmentation (Davenport & Ronanki, 2018; Shneiderman, 2020). Theoretically, this shift implies that authority in contemporary organizations is no longer located solely within hierarchical structures or formal roles but is distributed across socio-technical systems. As a result, legitimacy must be understood as co-produced through ongoing interactions between human cognition and algorithmic processes.

Third, the study contributes to the literature by adopting a multilevel perspective that integrates individual, organizational, and institutional dimensions of legitimacy. While prior research often examines legitimacy at a single level, this study demonstrates that legitimacy in AI-enabled organizations emerges from cross-level interactions. At the individual level, cognitive trust, perceived fairness, and interpretability shape how actors evaluate algorithmic decisions (Dwork et al., 2012). At the organizational level, governance structures, transparency mechanisms, and accountability systems influence how algorithmic outputs are framed and justified (Davenport & Ronanki, 2018). At the institutional level, broader regulatory, cultural, and normative frameworks define the boundaries of acceptable algorithmic authority (Meyer & Rowan, 1977; Yeung, 2018). By integrating these levels, the study contributes to a more holistic and dynamic understanding of legitimacy as an emergent, context-dependent phenomenon.

Fourth, the study enriches the growing literature on AI ethics and governance by linking technical attributes of AI systems—such as fairness, transparency, and explainability—to broader organizational and institutional processes. While prior research has extensively examined these attributes from a technical or normative perspective (Dwork et al., 2012; Wachter et al., 2017), this study embeds them within a legitimacy framework, thereby connecting micro-level system characteristics with macro-level perceptions of appropriateness and trustworthiness. This integration provides a theoretical bridge between algorithmic design and organizational outcomes, highlighting that ethical AI is not only a technical challenge but also a socio-organizational one.

Finally, the study contributes by reframing legitimacy as a dynamic and continuously negotiated construct in AI-enabled environments. Unlike traditional views that treat legitimacy as relatively stable once established, the proposed framework emphasizes its fluid and processual nature. The presence of opaque, adaptive, and evolving algorithmic systems introduces new sources of uncertainty, requiring ongoing justification and validation of decisions. This perspective aligns with contemporary shifts in organizational theory toward process-based and relational understandings of key constructs, and it underscores the need to conceptualize legitimacy as an ongoing accomplishment rather than a static attribute (Mittelstadt et al., 2016).

Taken together, these contributions not only extend existing theoretical boundaries but also open new avenues for integrating organizational theory with the rapidly evolving field of AI. By positioning legitimacy at the intersection of technology, human cognition, and institutional structures, the study provides a robust foundation for future theory-building efforts in AI-enabled organizational research.

7. PRACTICAL IMPLICATIONS

This study offers several practical implications for managers, policymakers, and system designers operating in AI-enabled organizational environments. As organizations increasingly rely on algorithmic systems to support or fully automate decision-making processes, the question of how to ensure decision legitimacy becomes not only a theoretical concern but also a critical managerial challenge.

First, the findings suggest that organizations should not assume that the technical accuracy or efficiency of AI systems is sufficient to guarantee acceptance of algorithmic decisions. Instead, legitimacy depends heavily on how these systems are perceived by organizational members. Therefore, managers should prioritize the development of transparent, explainable, and communicable AI systems. Providing clear rationales for algorithmic outputs can significantly enhance perceived fairness and reduce resistance to AI-driven decisions. This aligns with prior research emphasizing the importance of explainability and transparency in fostering trust in algorithmic systems (Wachter, Mittelstadt, & Russell, 2017; Shneiderman, 2020).

Second, organizations should actively design hybrid decision-making structures that integrate human judgment with algorithmic recommendations. Rather than fully delegating decisions to AI systems, organizations benefit from maintaining human oversight, particularly in contexts involving ethical ambiguity, high uncertainty, or high stakeholder impact. Human actors play a crucial role in interpreting, contextualizing, and validating algorithmic outputs. This hybrid governance approach not only enhances decision quality but also strengthens perceptions of procedural justice and legitimacy (Davenport & Ronanki, 2018).

Third, the study highlights the importance of developing organizational capabilities that support algorithmic literacy among employees and managers. Without a basic understanding of how AI systems generate outputs, organizational members may perceive algorithmic decisions as opaque or arbitrary, which can undermine trust and acceptance. Training programs that enhance AI literacy and foster critical engagement with algorithmic systems can help bridge this gap and enable more informed interactions between humans and AI systems.

Fourth, at the organizational level, firms should establish formal governance mechanisms to regulate the use of AI in decision-making processes. This includes defining accountability structures, setting ethical guidelines, and ensuring that responsibility for algorithmic outcomes is clearly assigned. Such mechanisms are essential for maintaining institutional legitimacy and ensuring that AI systems are perceived as fair, accountable, and aligned with organizational values (Yeung, 2018; Mittelstadt et al., 2016).

Fifth, the findings suggest that attention must be paid to the emotional and psychological responses of employees to algorithmic decision-making. The introduction of AI systems may generate uncertainty, perceived loss of autonomy, or concerns about surveillance and control. Organizations should therefore implement change management strategies that address these concerns proactively, emphasizing the supportive and augmentative role of AI rather than its substitutive potential. This can help mitigate resistance and foster a more positive acceptance of AI-enabled transformation.

Finally, policymakers and regulators should consider the broader implications of algorithmic authority in organizational contexts. As AI systems become increasingly embedded in decision-making processes, there is a growing need for regulatory frameworks that ensure transparency, fairness, and accountability. Such frameworks should not only focus on technical standards but also address the social and organizational dimensions of algorithmic governance.

Overall, the practical implications of this study underscore that successful AI integration is not solely a technological challenge but a socio-organizational one. Organizations that recognize the importance of legitimacy, human judgment, and multilevel governance will be better positioned to harness the benefits of AI while minimizing resistance and unintended consequences.

8. FUTURE RESEARCH AGENDA

Building on the proposed multilevel framework and theoretical propositions, this study opens several promising avenues for future research on decision legitimacy in AI-enabled organizations. In particular, future studies can extend, test, and refine the relationships articulated in the propositions by adopting diverse methodological approaches and contextual settings.

First, future research should empirically test the proposed relationships between algorithmic authority, human judgment, and perceived decision legitimacy. While this study offers a conceptual foundation, Proposition 1 and Proposition 2 suggest that legitimacy is shaped by the interaction between perceived algorithmic authority and human interpretive processes. Quantitative studies using survey-based designs could examine how perceptions of algorithmic explainability, transparency, and fairness influence legitimacy outcomes across different organizational roles and hierarchical levels. Experimental designs may also be useful in isolating causal effects of algorithmic versus human decision framing on legitimacy perceptions.

Second, longitudinal research designs are needed to better understand the dynamic and evolving nature of legitimacy proposed in Proposition 3. Since legitimacy is conceptualized as a continuously negotiated process rather than a static outcome, future studies should investigate how legitimacy perceptions change over time as employees gain experience with AI systems. Such studies could explore whether initial resistance to algorithmic authority diminishes with familiarity or whether sustained exposure leads to institutionalized acceptance or persistent skepticism.

Third, future research should further explore the multilevel mechanisms outlined in Proposition 4 by examining cross-level interactions between individual, organizational, and institutional factors. Multi-level modeling approaches could be used to assess how individual-level trust in AI interacts with organizational governance structures and institutional norms to shape legitimacy perceptions. Comparative studies across industries or cultural contexts would also be particularly valuable in identifying boundary conditions for the proposed framework.

Fourth, qualitative research approaches such as case studies and ethnographic investigations could deepen understanding of how human actors actively interpret and negotiate algorithmic authority in practice. Such approaches would be particularly useful for unpacking the micro-level cognitive and behavioral processes underlying Proposition 2, where human judgment is conceptualized as a moderator of algorithmic decision acceptance. In-depth organizational studies could reveal how legitimacy is co-constructed through everyday interactions between employees and AI systems.

Fifth, future research should examine boundary conditions that may weaken or strengthen the relationships proposed in the framework. For instance, Proposition 5 suggests that contextual factors such as task complexity, decision criticality, and ethical sensitivity may moderate the relationship between algorithmic authority and legitimacy. Investigating these contingencies would help refine the generalizability of the model and provide more nuanced insights into when algorithmic decision-making is more or less likely to be perceived as legitimate.

Finally, future studies should extend the framework to emerging technological contexts, such as generative AI, autonomous agents, and fully automated decision systems. These technologies may further complicate the balance between algorithmic authority and human judgment, potentially challenging existing assumptions about control, accountability, and legitimacy. Exploring these developments would allow researchers to continuously update and refine theoretical understandings of decision legitimacy in rapidly evolving digital environments.

Overall, the propositions developed in this study provide a structured foundation for future empirical, qualitative, and mixed-method research. By systematically investigating these propositions across different levels of analysis and contexts, scholars can further advance understanding of how legitimacy is constructed, maintained, and transformed in AI-enabled organizations.

9. CONCLUSION

This study develops a multilevel conceptual framework to explain how decision legitimacy is constructed, negotiated, and sustained in AI-enabled organizations through the interaction of algorithmic authority and human judgment. By integrating insights from legitimacy theory, AI governance, and organizational behavior, it demonstrates that legitimacy is not a static property of either technological systems or managerial authority, but rather an emergent outcome of ongoing socio-technical interactions across individual, organizational, and institutional levels. Within this view, the proposed theoretical propositions emphasize that algorithmic authority gains legitimacy not merely through technical performance, but through its interpretability, perceived fairness, and alignment with human values as enacted and reinforced by organizational actors.

By positioning human judgment as a critical moderating and sensemaking mechanism, the study challenges deterministic perspectives that conceptualize AI as a substitute for human decision-making. Instead, AI-enabled organizations are framed as hybrid systems in which legitimacy is co-produced through continuous negotiation between human actors and algorithmic systems. This perspective advances existing theory by reframing legitimacy as a dynamic, relational, and multilevel construct that evolves in parallel with technological change.

Overall, the framework contributes to a deeper theoretical understanding of how organizations can navigate the increasing complexity introduced by artificial intelligence while sustaining legitimacy in the eyes of stakeholders. It also provides a foundation for future empirical research to test and refine the proposed relationships, ultimately supporting the development of more accountable, transparent, and human-centered AI-enabled organizational systems.

Building on this foundation, the study concludes that the legitimacy of AI-enabled decision-making is not merely a function of technological efficacy but is fundamentally shaped by a dynamic interplay between algorithmic capabilities and human interpretive judgment across multiple organizational levels. From this perspective, legitimacy emerges as a continuously negotiated and reconstructed outcome within hybrid decision environments, rather than a fixed attribute of either technology or organizational processes. Accordingly, this perspective calls for a conceptual shift away from viewing AI as a discrete technical intervention toward understanding it as an integral and evolving component of organizational socio-technical systems, in which legitimacy is constantly redefined through ongoing interactions between human and algorithmic actors.

Future research should therefore examine how organizations—particularly internal auditing functions—reconstruct legitimacy during radical technological transformations, such as the implementation of advanced AI systems that challenge established domains of human expertise and autonomy (Elbardan et al., 2023). In addition, longitudinal studies are needed to better understand the long-term implications of AI integration on decision-making processes and organizational structures, especially given the rapid evolution of AI technologies and their dynamic effects on human–AI collaboration models (Guler et al., 2025; Singh et al., 2025). Within this scope, quantitative research could further explore the relationship between varying levels of AI literacy among managers and the perceived legitimacy of AI-generated recommendations across different organizational contexts (Passlack et al., 2025; Sabharwal et al., 2024). Relatedly, future studies may extend this inquiry by collecting detailed individual-level data on organizational resources and capabilities that shape these perceptions (Pinski et al., 2024).

Moreover, additional research is needed to investigate how organizations can leverage AI systems to foster continuous learning and adaptive decision-making processes, which are essential for sustaining long-term competitiveness (Shukla, 2024). At the team level, examining the influence of AI on decision dynamics—particularly trust formation, coordination mechanisms, and shared cognition—represents another critical avenue for future inquiry (Li & Tian, 2025). In this context, studies should also evaluate the design and effectiveness of training programs aimed at enhancing human–AI collaboration, as such interventions are crucial for setting realistic expectations and strengthening critical analytical skills in AI-supported environments (Eigner & Handler, 2024). Furthermore, developing frameworks that clearly delineate the roles and responsibilities of both human and AI agents across different business contexts is essential for understanding hybrid intelligence dynamics and ensuring that AI systems augment rather than undermine human decision-making processes (Söllner et al., 2025).

Finally, further empirical research is required to test the assumption that AI should primarily augment rather than automate human judgment, by systematically examining how AI can best support human decision-making without displacing human agency (Dwivedi et al., 2019).

REFERENCES

- 1) Ahmed, I. (2025). Navigating ethics and risk in artificial intelligence applications within information technology: A systematic review [Review of Navigating ethics and risk in artificial intelligence applications within information technology: A systematic review]. *American Journal of Advanced Technology and Engineering Solutions*, 1(1), 579–601. <https://doi.org/10.63125/590d7098> Akbarighatar, P. (2024). Operationalizing responsible AI principles through responsible AI capabilities. *AI and Ethics*. <https://doi.org/10.1007/s43681-024-00524-4>
- 2) Alhosani, K., & Alhashmi, S. M. (2024). Opportunities, challenges, and benefits of AI innovation in government services: a review [Review of Opportunities, challenges, and benefits of AI innovation in government services: a review]. *Discover Artificial Intelligence*, 4(1). Springer Nature. <https://doi.org/10.1007/s44163-024-00111-w>
- 3) Ananny, M., & Crawford, K. (2018). Seeing without knowing: Limitations of the transparency ideal and its application to algorithmic accountability. *New Media & Society*, 20(3), 973–989. <https://doi.org/10.1177/1461444816676645>
- 4) Asatiani, A., Malo, P., Nagbøl, P. R., Penttinen, E., Rinta-Kahila, T., & Salovaara, A. (2021). Sociotechnical Envelopment of Artificial Intelligence: An Approach to Organizational Deployment of Inscrutable Artificial Intelligence Systems. *Journal of the Association for Information Systems*, 22(2), 325–352. <https://doi.org/10.17705/1jais.00664>
- 5) Ashforth, B. E., & Gibbs, B. W. (1990). The double-edge of organizational legitimation. *Organization Science*, 1(2), 177–194. <https://doi.org/10.1287/orsc.1.2.177>
- 6) Badgami, R. K. (2025). AI as a Shadow Executive: Modelling Non-Human Influence in Organizational Decision-Making. *International Journal of Innovative Science and Research Technology (IJISRT)*, 1626–1626. <https://doi.org/10.38124/ijisrt/25dec1016>
- 7) Barocas, S., & Selbst, A. D. (2016). Big data's disparate impact. *California Law Review*, 104(3), 671–732.
- 8) Batool, A., Zowghi, D., & Bano, M. (2025). AI governance: a systematic literature review. *AI and Ethics*, 5(3), 3265–3279. <https://doi.org/10.1007/s43681-024-00653-w>
- 9) Beckert, J. (1999). Agency, entrepreneurs, and institutional change. *Organization Studies*, 20(5), 777–799. <https://doi.org/10.1177/0170840699205004>
- 10) Beer, D. (2017). The social power of algorithms. *Information, Communication & Society*, 20(1), 1–13. <https://doi.org/10.1080/1369118X.2016.1216147>
- Binns, R. (2018). Fairness in machine learning: Lessons from political philosophy. *Proceedings of Machine Learning Research*, 81, 1–11.
- 11) Binns, R., Veale, M., Van Kleek, M., & Shadbolt, N. (2018). 'It's reducing a human being to a percentage': Perceptions of justice in algorithmic decisions. *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems*, 1–14. <https://doi.org/10.1145/3173574.3173951>
- 12) Bitektine, A., & Haack, P. (2015). The “macro” and the “micro” of legitimacy: Toward a multilevel theory of the legitimacy process. *Academy of Management Review*, 40(1), 49–75. <https://doi.org/10.5465/amr.2013.0318>
- 13) Brynjolfsson, E., & McAfee, A. (2017). *Machine, platform, crowd*. W. W. Norton.
- 14) Burrell, J. (2016). How the machine ‘thinks’: Understanding opacity in machine learning algorithms. *Big Data & Society*, 3(1). <https://doi.org/10.1177/2053951715622512>

- 15) Burrell, J. (2016). How the machine 'thinks': Understanding opacity in machine learning algorithms. *Big Data & Society*, 3(1). <https://doi.org/10.1177/2053951715622512>
- 16) Cabiddu, F., Moi, L., Patriotta, G., & Allen, D. G. (2022). Why do users trust algorithms? A review and conceptualization of initial trust and trust over time [Review of *Why do users trust algorithms? A review and conceptualization of initial trust and trust over time*]. *European Management Journal*, 40(5), 685–706. Elsevier BV. <https://doi.org/10.1016/j.emj.2022.06.001>
- 17) Cahyani, R. R., & Musslifah, A. R. (2025). Balancing bytes and biases: A case study of AI adoption in academic human resource management. *Journal of Educational Management and Instruction (JEMIN)*, 5(2), 437–450. <https://doi.org/10.22515/jemin.v5i2.11679>
- 18) Castelo, N., Bos, M. W., & Lehmann, D. R. (2019). Task-dependent algorithm aversion. *Journal of Marketing Research*, 56(5), 809–825. <https://doi.org/10.1177/0022243719851788>
- 19) Chiu, Y.-T., Zhu, Y., & Corbett, J. (2021). In the hearts and minds of employees: A model of pre-adoptive appraisal toward artificial intelligence in organizations. *International Journal of Information Management*, 60, 102379–102379. <https://doi.org/10.1016/j.ijinfomgt.2021.102379>
- 20) Choung, H., David, P., & Seberger, J. S. (2023). A multilevel framework for AI governance. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2307.03198>
- 21) Colquitt, J. A., Conlon, D. E., Wesson, M. J., Porter, C. O., & Ng, K. Y. (2001). Justice at the millennium. *Journal of Applied Psychology*, 86(3), 425–445. <https://doi.org/10.1037/0021-9010.86.3.425>
- 22) Danaher, J., Hogan, M. J., Noone, C., Kennedy, R., Behan, A., De Paor, A., ... & Shankar, K. (2017). Algorithmic governance. *Big Data & Society*, 4(2). <https://doi.org/10.1177/2053951717718855>
- 23) Davenport, T. H., & Kirby, J. (2016). *Only humans need apply*. Harper Business.
- 24) Davenport, T. H., & Ronanki, R. (2018). Artificial intelligence for the real world. *Harvard Business Review*, 96(1), 108–116.
- 25) De Cremer, D., & Kaspárov, G. (2021). *The future of work: AI and human collaboration in organizations*. Harvard Business Review Press.
- 26) Deephouse, D. L., Bundy, J., Tost, L. P., & Suchman, M. C. (2017). Organizational legitimacy. *Academy of Management Annals*, 11(1), 451–478. <https://doi.org/10.5465/annals.2015.0101>
- 27) Demirsoy, S. (2025). Business management in algorithmic environments: Redefining accountability, governance, and performance. Zenodo (CERN European Organization for Nuclear Research). <https://doi.org/10.5281/zenodo.18501845>
- 28) Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion. *Journal of Experimental Psychology: General*, 144(1), 114–126. <https://doi.org/10.1037/xge0000033>
- 29) Dima, J., Gilbert, M., Dextras-Gauthier, J., & Giraud, L. (2024). The effects of artificial intelligence on human resource activities and the roles of the human resource triad: opportunities and challenges. *Frontiers in Psychology*, 15. <https://doi.org/10.3389/fpsyg.2024.1360401>
- 30) DiMaggio, P. J., & Powell, W. W. (1983). The iron cage revisited. *American Sociological Review*, 48(2), 147–160.
- 31) Dwivedi, Y. K., Hughes, L., Ismagilova, E., Aarts, G., Coombs, C., Crick, T., Duan, Y., Dwivedi, R., Edwards, J. S., Eirug, A., Galanos, V., Ilavarasan, P. V., Janssen, M., Jones, P., Kar, A. K., Kizgin, H., Kronemann, B., Lal, B., Lucini, B., ... Williams, M. D. (2019). Artificial Intelligence (AI): Multidisciplinary perspectives on emerging challenges, opportunities, and agenda for research, practice and policy. *International Journal of Information Management*, 57, 101994–101994. <https://doi.org/10.1016/j.ijinfomgt.2019.08.002>
- 32) Dwork, C., Hardt, M., Pitassi, T., Reingold, O., & Zemel, R. (2012). Fairness through awareness. *Proceedings of the 3rd Innovations in Theoretical Computer Science Conference*. <https://doi.org/10.1145/2090236.2090255>
- 33) Eigner, E., & Handler, T. (2024). Determinants of LLM-assisted Decision-Making. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2402.17385>
- 34) Elbardan, H., Nordberg, D., & Sinha, V. K. (2023). Reconstructing legitimacy of internal auditing during ERP implementations: two contrasting cases. *Journal of Accounting Literature*, 47(5), 184–210. <https://doi.org/10.1108/jal-01-2023-0001>
- 35) Enang, I., Mukala, P., & Nkerekwem, U. (2026). Authentic Intelligence in Digital Strategy Systems: A Socio-Technical Analysis of Human-Accountable Decision Governance. *Systems*, 14(3), 259–259. <https://doi.org/10.3390/systems14030259>
- 36) Enholm, I. M., Papagiannidis, E., Mikalef, P., & Krogstie, J. (2021). Artificial Intelligence and Business Value: a Literature Review [Review of *Artificial Intelligence and Business Value: a Literature Review*]. *Information Systems Frontiers*, 24(5), 1709–1734. Springer Science+Business Media. <https://doi.org/10.1007/s10796-021-10186-w>
- 37) Esch, P. van. (2026). From agentic AI to AI-orchestrated organizations: Understanding the next surge in artificial intelligence. *Business Horizons*. <https://doi.org/10.1016/j.bushor.2026.03.003>
- 38) Faraj, S., Pachidi, S., & Sayegh, K. (2018). Working and organizing in the age of algorithms. *Information and Organization*, 28(1), 62–70. <https://doi.org/10.1016/j.infoandorg.2018.02.005>
- 39) Gharaibeh, O. K. K. (2026a). AI As The Invisible Manager: How Artificial Intelligence Is Reshaping Managerial Decision-Making. Zenodo (CERN European Organization for Nuclear Research). <https://doi.org/10.5281/zenodo.18411174>
- 40) Gharaibeh, O. K. K. (2026b). AI As The Invisible Manager: How Artificial Intelligence Is Reshaping Managerial Decision-Making. Zenodo (CERN European Organization for Nuclear Research). <https://doi.org/10.5281/zenodo.18411175>
- 41) Gillespie, T. (2014). The relevance of algorithms. In T. Gillespie et al. (Eds.), *Media technologies*. MIT Press.
- 42) Gladden, M. E., Fortuna, P., & Modliński, A. (2022). The Empowerment of Artificial Intelligence in Post-Digital Organizations: Exploring Human Interactions with Supervisory AI. *Human Technology*, 18(2), 98–121. <https://doi.org/10.14254/1795-6889.2022.18-2.2>

- 43) Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence. *Academy of Management Review*, 45(2), 1–24. <https://doi.org/10.5465/amr.2018.0196>
- 44) Greenberg, J. (1990). Organizational justice. *Journal of Management*, 16(2), 399–432.
- 45) Guler, N., Cahalane, M., Kirshner, S. N., & Vidgen, R. (2025). The Role of Roles: When are LLMs Behavioural in Information Systems Decision-Making. *AJIS. Australasian Journal of Information Systems/AJIS. Australian Journal of Information Systems/Australian Journal of Information Systems*, 29. <https://doi.org/10.3127/ajis.v29.5573>
- 46) Haack, P., Pfarrer, M. D., & Scherer, A. G. (2014). Legitimacy-as-feeling. *Academy of Management Journal*, 57(3), 634–660.
- 47) Haesevoets, T., Cremer, D. D., Dierckx, K., & Hiel, A. V. (2021). Human-machine collaboration in managerial decision making. *Computers in Human Behavior*, 119, 106730–106730. <https://doi.org/10.1016/j.chb.2021.106730>
- 48) Highhouse, S. (2008). Stubborn reliance on intuition. *Industrial and Organizational Psychology*, 1(3), 333–342.
- 49) Hillebrand, L., Raisch, S., & Schad, J. (2025). Managing with Artificial Intelligence: An Integrative Framework. *Academy of Management Annals*, 19(1), 343–375. <https://doi.org/10.5465/annals.2022.0072>
- 50) Hossain, S., Fernando, M., & Akter, S. (2025a). Digital Leadership: Towards a Dynamic Managerial Capability Perspective of Artificial Intelligence-Driven Leader Capabilities. *Journal of Leadership & Organizational Studies*, 32(2), 189–208. <https://doi.org/10.1177/15480518251319624>
- 51) Hossain, S., Fernando, M., & Akter, S. (2025b). The influence of artificial intelligence-driven capabilities on responsible leadership: A future research agenda. *Journal of Management & Organization*, 31(5), 2360–2384. <https://doi.org/10.1017/jmo.2025.10010>
- 52) Huang, B., & Niyomsilp, E. (2025). The impact of artificial intelligence on organizational decision-making processes. *Edelweiss Applied Science and Technology*, 9(4), 794–808. <https://doi.org/10.55214/25768484.v9i4.6081>
- 53) Kahneman, D. (2011). *Thinking, fast and slow*. Farrar, Straus and Giroux.
- 54) Kahneman, D. (2011). *Thinking, fast and slow*. Farrar, Straus and Giroux.
- 55) Kellogg, K. C., Valentine, M. A., & Christin, A. (2020). Algorithms at work. *Academy of Management Annals*, 14(1), 366–410. <https://doi.org/10.5465/annals.2018.0174>
- 56) 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>
- 57) Koponen, J., Julkunen, S., Laajalahti, A., Turunen, M., & Spitzberg, B. H. (2023). Work Characteristics Needed by Middle Managers When Leading AI-Integrated Service Teams. *Journal of Service Research*. <https://doi.org/10.1177/10946705231220462>
- 58) Kroll, J. A., Huey, J., Barocas, S., Felten, E. W., Reidenberg, J. R., Robinson, D. G., & Yu, H. (2017). Accountable algorithms. *University of Pennsylvania Law Review*, 165(3), 633–705.
- 59) Lanz, L., Briker, R., & Gerpott, F. H. (2023). Employees Adhere More to Unethical Instructions from Human Than AI Supervisors: Complementing Experimental Evidence with Machine Learning. *Journal of Business Ethics*, 189(3), 625–646. <https://doi.org/10.1007/s10551-023-05393-1>
- 60) 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
- 61) Lee, M. K. (2018). Understanding perception of algorithmic decisions. *Proceedings of the ACM on Human-Computer Interaction*, 2(CSCW), 1–26. <https://doi.org/10.1145/3274429>
- 62) Leoni, L., Gueli, G., Ardolino, M., Panizzon, M., & Gupta, S. (2024). AI-empowered KM processes for decision-making: empirical evidence from worldwide organisations. *Journal of Knowledge Management*, 28(11), 320–347. <https://doi.org/10.1108/jkm-03-2024-0262>
- 63) Li, H., & Tian, F. (2025). Advancing Decision-Making through AI-Human Collaboration: A Systematic Review and Conceptual Framework [Review of *Advancing Decision-Making through AI-Human Collaboration: A Systematic Review and Conceptual Framework*]. *Research Square (Research Square)*. Research Square (United States). <https://doi.org/10.21203/rs.3.rs-6885768/v1>
- 64) Lind, E. A., & Tyler, T. R. (1988). *The social psychology of procedural justice*. Springer.
- 65) Logg, J. M., Minson, J. A., & Moore, D. A. (2019). Algorithm appreciation. *Organizational Behavior and Human Decision Processes*, 151, 90–103. <https://doi.org/10.1016/j.obhdp.2019.01.002>
- 66) Madanchian, M., & Taherdoost, H. (2025). Ethical theories, governance models, and strategic frameworks for responsible AI adoption and organizational success. *Frontiers in Artificial Intelligence*, 8, 1619029–1619029. <https://doi.org/10.3389/frai.2025.1619029>
- 67) Malik, A., Budhwar, P., Patel, C., & Srikanth, N. R. (2020). May the bots be with you! Delivering HR cost-effectiveness and individualised employee experiences in an MNE. *The International Journal of Human Resource Management*, 33(6), 1148–1178. <https://doi.org/10.1080/09585192.2020.1859582>
- 68) March, J. G., & Olsen, J. P. (1989). *Rediscovering institutions: The organizational basis of politics*. Free Press.
- 69) Martin, K., & Waldman, A. E. (2022). Are Algorithmic Decisions Legitimate? The Effect of Process and Outcomes on Perceptions of Legitimacy of AI Decisions. *Journal of Business Ethics*, 183(3), 653–670. <https://doi.org/10.1007/s10551-021-05032-7>
- 70) Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *Academy of Management Review*, 20(3), 709–734.
- 71) Meyer, J. W., & Rowan, B. (1977). Institutionalized organizations. *American Journal of Sociology*, 83(2), 340–363.
- 72) Meyer, J. W., & Rowan, B. (1977). Institutionalized organizations: Formal structure as myth and ceremony. *American Journal of Sociology*, 83(2), 340–363. <https://doi.org/10.1086/226550>
- 73) Mittelstadt, B. D., Allo, P., Taddeo, M., Wachter, S., & Floridi, L. (2016). The ethics of algorithms. *Big Data & Society*, 3(2).

- 74) 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>
- 75) Mohammed, A., Bambale, S. A., & Abdullahi, U. (2025). AI-Driven Decision-Making in Leadership: Balancing Automation and Human Judgment for Optimal Organizational Performance. *Journal of Economics and Administrative Sciences*, 31(150), 129–141. <https://doi.org/10.33095/hv149002>
- 76) Noble, S. U. (2018). *Algorithms of oppression*. NYU Press.
- 77) O'Neil, C. (2016). *Weapons of math destruction*. Crown.
- 78) O'Neil, C. (2016). *Weapons of math destruction*. Crown.
- 79) Oladele, I., Orelaja, A., & Akinwande, O. T. (2024). Ethical Implications and Governance of Artificial Intelligence in Business Decisions: A Deep Dive into the Ethical Challenges and Governance Issues Surrounding the Use of Artificial Intelligence in Making Critical Business Decisions. *International Journal of Latest Technology in Engineering Management & Applied Science*, 48–56. <https://doi.org/10.51583/ijltemas.2024.130207>
- 80) Ossa, L. A., Lorenzini, G., Milford, S. R., Shaw, D., Elger, B. S., & Rost, M. (2024). Integrating ethics in AI development: a qualitative study. *BMC Medical Ethics*, 25(1). <https://doi.org/10.1186/s12910-023-01000-0>
- 81) Papagiannidis, E., Mikalef, P., & Conboy, K. (2025). Responsible artificial intelligence governance: A review and research framework. *The Journal of Strategic Information Systems*, 34(2), 101885–101885. <https://doi.org/10.1016/j.jsis.2024.101885>
- 82) Pasquale, F. (2015). *The black box society*. Harvard University Press.
- 83) Passlack, N., Hammerschmidt, T., & Posegga, O. (2025). With Great Power Comes Great Responsibility: What Shapes AI Literacy for Responsible Interactions of Knowledge Workers With AI? *Information Systems Frontiers*. <https://doi.org/10.1007/s10796-025-10648-5>
- 84) Pinski, M., Hofmann, T., & Benlian, A. (2024). AI Literacy for the top management: An upper echelons perspective on corporate AI orientation and implementation ability. *Electronic Markets*, 34(1). <https://doi.org/10.1007/s12525-024-00707-1>
- 85) Pöhler, L., Diepold, K., & Wallach, W. (2024). A Practical Multilevel Governance Framework for Autonomous and Intelligent Systems. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2404.13719>
- 86) Qeybi, N., & Bösch, S. (2025). Integrating perspectives on democratization in industry via multi-agent systems: insights from a firm-based case study. *Human-Intelligent Systems Integration*. <https://doi.org/10.1007/s42454-025-00077-9>
- 87) Rai, A. (2020). Explainable AI: From black box to glass box. *Journal of the Academy of Marketing Science*, 48, 137–141. <https://doi.org/10.1007/s11747-019-00710-5>
- 88) 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>
- 89) Raisch, S., & Krakowski, S. (2021). Artificial intelligence and management. *Academy of Management Journal*, 64(1), 192–210.
- 90) 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>
- 91) Reddy, P. S., Reddy, M. R., Hansraj, B. D., Archana, G., Kanala, B., & Radhika, Ch. (2026a). How Machine Learning Reshapes Strategic Choices in Organizations- A Bibliometric Analysis. *Zenodo (CERN European Organization for Nuclear Research)*. <https://doi.org/10.5281/zenodo.18976157>
- 92) Reddy, P. S., Reddy, M. R., Hansraj, B. D., Archana, G., Kanala, B., & Radhika, Ch. (2026b). How Machine Learning Reshapes Strategic Choices in Organizations- A Bibliometric Analysis. *Journal of Economics Finance and Management Studies*, 9(3), 1289–1298. <https://doi.org/10.47191/jefms/v9-i3-13>
- 93) Sabharwal, R., Miah, S. J., Wamba, S. F., & Cook, P. (2024). Extending application of explainable artificial intelligence for managers in financial organizations. *Annals of Operations Research*. <https://doi.org/10.1007/s10479-024-05825-9>
- 94) Schmidt, J.-H., Bartsch, S. C., Adam, M., & Benlian, A. (2025). Elevating Developers' Accountability Awareness in AI Systems Development. *Business & Information Systems Engineering*, 67(1), 109–135. <https://doi.org/10.1007/s12599-024-00914-2>
- 95) Scott, W. R. (2014). *Institutions and organizations*. Sage.
- 96) Sharabati, A. A., Rehman, S. U., Malik, M., Sabra, S., Al-Sager, M., & Allahham, M. (2024). Is AI biased? evidence from FinTech-based innovation in supply chain management companies? *International Journal of Data and Network Science*, 8(3), 1839–1852. <https://doi.org/10.5267/j.ijdns.2024.2.005>
- 97) Shi, H. (2026). Artificial Intelligence-Embedded Restructuring of Enterprise Decision Authority: A Re-examination from the Perspective of Management Theory. *Economics & Business Management*, 5(1), 1–1. <https://doi.org/10.63313/ebm.9154>
- 98) Shneiderman, B. (2020). Human-centered artificial intelligence: Reliable, safe & trustworthy. *International Journal of Human–Computer Interaction*, 36(6), 495–504. <https://doi.org/10.1080/10447318.2020.1741118>
- 99) Shrestha, Y. R., Ben-Menahem, S. M., & Krogh, G. von. (2019). Organizational Decision-Making Structures in the Age of Artificial Intelligence. *California Management Review*, 61(4), 66–83. <https://doi.org/10.1177/0008125619862257>
- 100) Shukla, Mr. M. (2024). The Impact of AI on Improving the Efficiency and Accuracy of Managerial Decisions. *International Journal for Research in Applied Science and Engineering Technology*, 12(7), 830–842. <https://doi.org/10.22214/ijraset.2024.63652>
- 101) Simon, H. A. (1997). *Administrative behavior* (4th ed.). Free Press.
- 102) Singh, K., Bojilov, M., & Best, P. (2025). Implementing AI in auditing in organizations. *Journal of Accounting and Management Information Systems*, 24(3). <https://doi.org/10.24818/jamis.2025.03003>

- 103) Söllner, M., Arnold, T., Benlian, A., Bretschneider, U., Knight, C. L., Ohly, S., Rudkowski, L., Schreiber, G., & Wendt, D. H. (2025). ChatGPT and Beyond: Exploring the Responsible Use of Generative AI in the Workplace. *Business & Information Systems Engineering*. <https://doi.org/10.1007/s12599-025-00932-8>
- 104) Stahl, B. C., Antoniou, J., Ryan, M., Macnish, K., & Jiya, T. (2021). Organisational responses to the ethical issues of artificial intelligence. *AI & Society*, 37(1), 23–37. <https://doi.org/10.1007/s00146-021-01148-6>
- 105) Suchman, M. C. (1995). Managing legitimacy. *Academy of Management Review*, 20(3), 571–610.
- 106) Suchman, M. C. (1995). Managing legitimacy. *Academy of Management Review*, 20(3), 571–610. <https://doi.org/10.2307/258788>
- 107) Suchman, M. C. (1995). Managing legitimacy: Strategic and institutional approaches. *Academy of Management Review*, 20(3), 571–610 <https://doi.org/10.5465/amr.1995.9508080331>
- 108) Suchman, M. C. (1995). Managing legitimacy: Strategic and institutional approaches. *Academy of Management Review*, 20(3), 571–610 <https://doi.org/10.5465/amr.1995.9508080331>
- 109) Suddaby, R., Bitektine, A., & Haack, P. (2017). Legitimacy. *Academy of Management Annals*, 11(1), 451–478.
- 110) Sudeeptha, I., Mueller, W., Leyer, M., Richter, A., & Nolte, F. (2024). Use Cases for Prospective Sensemaking of Human-AI Collaboration. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2408.10812>
- 111) Tost, L. P. (2011). An integrative model of legitimacy judgments. *Academy of Management Review*, 36(4), 686–710.
- 112) Trist, E. L., & Bamforth, K. W. (1951). Some social and psychological consequences of the longwall method of coal-getting. *Human Relations*, 4(1), 3–38. <https://doi.org/10.1177/001872675100400101>
- 113) Trust and responsibility in AI-driven leadership: accountability for ai decisions. (2026). *Scientific Papers of Silesian University of Technology Organization and Management Series*, 2026(239). <https://doi.org/10.29119/1641-3466.2026.239.1>
- 114) Tyler, T. R. (2006). *Why people obey the law*. Princeton University Press.
- 115) Tyler, T. R., & Blader, S. (2003). The group engagement model. *Personality and Social Psychology Review*, 7(4), 349–361.
- 116) Upase, Dr. P. B. (2026). Reframing Organizational Decision-Making in the Age of Artificial Intelligence: A Conceptual Review of Human–AI Augmentation. *INTERANTIONAL JOURNAL OF SCIENTIFIC RESEARCH IN ENGINEERING AND MANAGEMENT*, 10(2), 1–9. <https://doi.org/10.55041/ijrsrem.ibfe071>
- 117) Übellacker, T. (2025). Making Sense of AI Limitations: How Individual Perceptions Shape Organizational Readiness for AI Adoption. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2502.15870>
- 118) Wachter, S., Mittelstadt, B., & Russell, C. (2017). Counterfactual explanations without opening the black box. *Harvard Journal of Law & Technology*, 31(2), 841–887.
- 119) Weber, M., Engert, M., Schaffer, N., Weking, J., & Krcmar, H. (2022). Organizational Capabilities for AI Implementation—Coping with Inscrutability and Data Dependency in AI. *Information Systems Frontiers*, 25(4), 1549–1569. <https://doi.org/10.1007/s10796-022-10297-y>
- 120) Westover, J. (2025). AI-Augmented Decision Rights: Redesigning Authority in Human-Machine Organizations. *Human Capital Leadership*, 27(1). <https://doi.org/10.70175/hclreview.2020.27.1.7>
- 121) Wójcik, M. (2024). Algorithmic discrimination in the era of artificial intelligence: challenges of sustainable human resource management. *Edukacja Ekonomistów i Menedżerów*, 69(3). <https://doi.org/10.33119/eeim.2024.69.6>
- 122) Wójcik, M. (2025). AI meets HR – systematic review of implementation and operational barriers. *Scientific Papers of Silesian University of Technology Organization and Management Series*, 2025(224), 643–660. <https://doi.org/10.29119/1641-3466.2025.224.35>
- 123) Xiong, M., Wang, Y., & Olya, H. (2025). ‘Real-ising’ the Benefits of Responsible Artificial Intelligence (AI) Management: An Interpretive Study (pp. 81–114). <https://doi.org/10.1108/978-1-80262-963-720251009>
- 124) Yeung, K. (2018). Algorithmic regulation: A critical interrogation. *Regulation & Governance*, 12(4), 505–523. <https://doi.org/10.1111/rego.12158>
- 125) Yiğitcanlar, T., Agdas, D., & Degirmenci, K. (2022). Artificial intelligence in local governments: perceptions of city managers on prospects, constraints and choices. *AI & Society*, 38(3), 1135–1150. <https://doi.org/10.1007/s00146-022-01450-x>
- 126) Zerilli, J., Knott, A., Maclaurin, J., & Gavaghan, C. (2019). Transparency in algorithmic decision-making. *Philosophy & Technology*, 32(4), 661–678.
- 127) Zuboff, S. (2019). *The age of surveillance capitalism*. PublicAffairs.
- 128) Zuboff, S. (2019). *The age of surveillance capitalism: The fight for a human future at the new frontier of power*. PublicAffairs.