

The Emergence of Algorithmic Leadership: A Grounded Theory Exploration of Managerial Authority in AI-Augmented Organizations

Hamideh Soltani¹, Pouria Afshari^{2*}

1. Department of Educational Management, Shi.C., Islamic Azad University, Shiraz, Iran

2. Department of Management, Mar.C., Islamic Azad University, Marvdasht, Iran

Abstract

This study aimed to develop a grounded theory explaining how managerial authority is reconstructed when artificial intelligence systems become embedded in decision-making, monitoring, coordination, and performance-management processes in AI-augmented organizations. This qualitative study used a grounded theory design to explore the emergence of algorithmic leadership in organizations located in Tehran. Data were collected through semi-structured interviews with 26 participants, including senior managers, middle managers, HR specialists, data and AI managers, and team leaders working in AI-augmented organizations. Participants were selected through purposive and theoretical sampling, and interviews continued until theoretical saturation was achieved. The interviews focused on experiences of AI-supported decision-making, changes in managerial roles, perceived legitimacy of algorithmic recommendations, accountability, employee trust, and leadership adaptation. Interviews were audio-recorded with consent, transcribed verbatim, and analyzed using open, axial, and selective coding. NVivo software was used to organize codes, compare categories, retrieve quotations, and develop the final grounded-theory model. The analysis generated five main categories: algorithmic delegation of managerial judgment, redistribution of authority between humans and AI systems, datafied visibility and control, leadership sensemaking and algorithmic literacy, and governance of algorithmic legitimacy. The core category was identified as calibrated algorithmic authority, referring to the dynamic process through which managers gradually transfer selected managerial functions to AI systems while retaining responsibility for interpretation, ethical judgment, exception handling, and employee meaning-making. Participants described AI as neither a neutral tool nor a full substitute for leadership, but as a new authority-bearing actor that reshapes how decisions are justified, contested, and implemented. The study suggests that algorithmic leadership emerges when AI systems become active participants in organizational authority rather than merely technical decision-support tools. In AI-augmented organizations, effective leadership depends on managers' ability to calibrate the relationship between algorithmic recommendations and human judgment. The proposed grounded theory highlights that algorithmic leadership requires technical literacy, ethical reflexivity, procedural transparency, accountability structures, and communicative competence. Organizations that treat AI as an unquestioned authority risk weakening trust and human agency, whereas organizations that govern AI as a transparent and contestable partner may strengthen decision quality, coordination, and adaptive leadership.

Keywords: Algorithmic leadership, algorithmic management, artificial intelligence, managerial authority, AI-augmented organizations, grounded theory, human-AI collaboration, Tehran.

✉ *CORRESPONDING AUTHOR'S EMAIL

pouria.afshari@iau.ac.ir

📅 ARTICLE HISTORY

Received date: 12 May 2024

Revised date: 13 June 2024

Accepted date: 25 June 2024

Published date: 01 July 2024

📖 HOW TO CITE

Soltani, H. & Afshari, P. (2024). The Emergence of Algorithmic Leadership: A Grounded Theory Exploration of Managerial Authority in AI-Augmented Organizations. *Management Systems Evolution and Leadership Studies*, 1(1), 34-45.

1. Introduction

Artificial intelligence has moved from the margins of organizational experimentation to the center of managerial practice. In contemporary organizations, AI systems are increasingly used to forecast demand, evaluate employee performance, allocate tasks, screen applicants, monitor workflow, recommend strategic



Copyright: © 2024 by the authors. Published under the terms and conditions of Creative Commons Attribution-NonCommercial 4.0 International (CC BY-NC 4.0) License.

options, detect anomalies, and support managerial decisions. This transition has produced a new organizational condition in which leadership is no longer exercised exclusively through human hierarchy, interpersonal influence, or bureaucratic procedures. Instead, authority is increasingly mediated by algorithmic systems that collect data, classify behavior, generate recommendations, and shape managerial attention. Algorithmic management has been described as the use of software, sometimes including AI, to automate or partially automate functions traditionally performed by human managers, including monitoring, allocation, evaluation, and supervision (Parent-Rocheleau & Parker, 2022). The International Labour Organization similarly defines algorithmic management as systems that use tracked data to organize, assign, monitor, supervise, and evaluate work. Within this context, the concept of algorithmic leadership becomes necessary because AI does not merely assist management; it changes the very basis upon which managerial authority is produced, justified, accepted, and contested.

Leadership studies have traditionally focused on human traits, behaviors, relationships, identities, influence processes, and institutional contexts. Transformational, distributed, adaptive, and digital leadership perspectives have all expanded the field beyond command-and-control assumptions, yet they still commonly treat technology as an environmental condition or a managerial resource rather than an authority-bearing actor. Digital leadership research has shown that leaders must guide organizations through technological transformation, ambiguity, and accelerated information flows (Avolio et al., 2014; Cortellazzo et al., 2019). However, AI intensifies this transformation because algorithmic systems increasingly perform tasks once considered central to managerial agency. They do not only provide information; they rank options, structure attention, define deviations, predict risk, and recommend interventions. As a result, managerial authority becomes hybridized. It is neither fully human nor fully automated, but co-produced through the interaction between managers, employees, AI systems, data infrastructures, and organizational rules.

The rise of algorithmic leadership is closely related to the broader literature on algorithmic management. Parent-Rocheleau and Parker (2022) identify several managerial functions that algorithms are able to perform, including monitoring, goal setting, performance management, scheduling, compensation, and termination. This body of work suggests that algorithmic systems are not passive tools but active organizers of work. Kellogg, Valentine, and Christin (2020) conceptualize algorithms at work as a new contested terrain of control, emphasizing that algorithmic systems reshape power relations through mechanisms that render work more visible, direct behavior, and structure reward or sanction. Similarly, Faraj, Pachidi, and Sayegh (2018) argue that learning algorithms transform expertise, occupational boundaries, and forms of coordination and control. These studies are highly relevant to leadership because they show that AI changes not only the technical execution of managerial tasks but also the social architecture of authority.

The managerial consequences of AI cannot be reduced to a simple opposition between automation and augmentation. Raisch and Krakowski (2021) argue that in management contexts, automation and augmentation are deeply interdependent rather than neatly separable. This insight is central to the present study. In many AI-augmented organizations, managers do not simply choose between making decisions themselves or delegating decisions to machines. Instead, they work within a shifting continuum in which AI may detect patterns, generate recommendations, prioritize issues, predict outcomes, or enforce rules, while human managers interpret these outputs, negotiate exceptions, communicate decisions, and absorb accountability. Thus, algorithmic leadership

should be understood as a dynamic configuration of human and machine agency rather than as the replacement of managers by AI.

Human–AI collaboration research further supports this perspective. Jarrahi (2018) argues that humans and AI contribute different strengths to organizational decision-making, especially in conditions marked by uncertainty, complexity, and equivocality. AI systems may be superior in processing large volumes of structured data, detecting patterns, and producing consistent outputs, whereas human managers remain essential for contextual interpretation, ethical judgment, empathy, political awareness, and sensemaking. Shrestha, Ben-Menahem, and von Krogh (2019) identify multiple structures through which human and AI decisions may be combined, including delegation, sequential decision-making, and aggregation. These frameworks suggest that the central question is not whether AI should lead, but how authority should be distributed, governed, and legitimized when managerial judgment is increasingly intertwined with algorithmic outputs.

Despite the growing importance of AI in organizations, the specific concept of managerial authority remains underdeveloped in much of the literature on AI and leadership. Existing studies often focus on decision quality, technological adoption, ethical risk, employee surveillance, or human trust in AI. These are important issues, yet they do not fully explain how managers experience the transformation of their authority when AI systems intervene in daily leadership functions. Authority is not merely the formal right to decide; it is also the socially recognized capacity to justify decisions, allocate responsibility, resolve ambiguity, and secure compliance. In AI-augmented organizations, authority may become ambiguous because the source of a decision is no longer always clear. When a performance score is generated by an AI system, when a promotion shortlist is produced through predictive analytics, or when an operational dashboard flags underperformance, the manager must decide whether to accept, override, reinterpret, or challenge the algorithmic output. The decision may appear data-driven, but the accountability remains organizational and human.

This ambiguity is intensified by issues of trust, explainability, and fairness. Glikson and Woolley (2020) show that trust in AI depends on factors such as the representation of AI, perceived intelligence, reliability, and task characteristics. In managerial settings, trust is not only a cognitive evaluation of system accuracy; it is also an organizational judgment about legitimacy. Employees may accept AI-supported decisions when they perceive them as consistent, transparent, and procedurally fair, but may resist them when they appear opaque, dehumanizing, or disconnected from contextual realities. Binns et al. (2018) demonstrate that perceptions of justice in algorithmic decisions involve concerns such as arbitrariness, generalization, and dignity. Therefore, algorithmic leadership requires more than technical implementation. It requires managers to create interpretive bridges between algorithmic outputs and human expectations of fairness, recognition, and voice.

AI also alters the role of middle managers. Traditionally, middle managers translate strategy into action, monitor performance, coordinate teams, interpret organizational expectations, and mediate between senior leadership and employees. In AI-augmented organizations, many of these functions are partially automated or reshaped by dashboards, predictive indicators, workflow systems, and performance analytics. Lee et al. (2015) showed that algorithmic and data-driven management can substantially influence worker experience and perceptions of managerial control. Tambe, Cappelli, and Yakubovich (2019) similarly warn that AI applications in HRM raise challenges related to complexity, small data, accountability, fairness, and employee reactions. These issues suggest that managers may increasingly become interpreters and governors of AI systems rather than sole originators of managerial decisions.

The need for qualitative inquiry is particularly strong because the emergence of algorithmic leadership is not only a technological phenomenon but also a lived organizational experience. Quantitative studies can measure adoption, trust, productivity, or performance outcomes, but they may not adequately capture how managers make sense of altered authority, how they negotiate accountability, or how they respond to employee concerns when AI systems influence managerial decisions. A grounded theory approach is appropriate because algorithmic leadership remains an emerging phenomenon and requires theory generation from participants' experiences rather than mere testing of predefined constructs. Grounded theory allows the researcher to examine processes, categories, conditions, interactions, and consequences as they emerge from empirical data (Charmaz, 2014; Corbin & Strauss, 2015).

The context of Tehran provides an important setting for this investigation. Organizations in Tehran increasingly operate in digitally intensive sectors such as fintech, e-commerce, telecommunications, banking, insurance, logistics, and platform-based services. These organizations face competitive pressure to adopt AI tools while also navigating institutional uncertainty, regulatory ambiguity, cultural expectations of managerial responsibility, and uneven levels of employee trust in automated decision-making. The study of algorithmic leadership in this context can extend existing theory by showing how managerial authority is reconstructed in organizations where AI adoption is practical and growing, but where governance norms, transparency practices, and human–AI collaboration routines may still be developing.

Accordingly, this study seeks to develop a grounded theory of algorithmic leadership by exploring how managers and organizational specialists experience the redistribution of authority in AI-augmented organizations. Rather than treating AI as either a threat to leadership or a neutral managerial tool, the study conceptualizes AI as part of a sociotechnical authority system. The central research question is: How does managerial authority emerge, shift, and become legitimized when AI systems participate in managerial decision-making and organizational control? The aim of this study is to develop a grounded theory explaining the emergence of algorithmic leadership and the reconstruction of managerial authority in AI-augmented organizations.

2. Methods and Materials

This study used a qualitative grounded theory design to explore the emergence of algorithmic leadership and the reconstruction of managerial authority in AI-augmented organizations. Grounded theory was selected because the phenomenon under investigation is still conceptually developing and requires an inductive explanation of processes, interactions, meanings, and consequences. The study followed the logic of constant comparison, theoretical sampling, simultaneous data collection and analysis, and progressive category development. The methodological orientation was constructivist grounded theory, which assumes that theory is generated through interaction between participants' meanings, organizational contexts, and the researcher's analytical interpretation (Charmaz, 2014). This design was suitable because algorithmic leadership is not simply a technical condition but a socially negotiated process through which managers, employees, and AI systems participate in the production of authority.

The participants were selected from AI-augmented organizations located in Tehran. For the purpose of this study, AI-augmented organizations were defined as organizations that used AI-based or algorithmic systems in at least one managerial function, such as performance monitoring, task allocation, recruitment screening, customer-

service analytics, risk scoring, demand forecasting, decision dashboards, employee evaluation, workflow optimization, or managerial reporting. The final sample included 26 participants. The participants consisted of 6 senior managers, 7 middle managers, 5 HR and people-analytics specialists, 4 AI or data-product managers, 2 compliance or risk managers, and 2 team leaders who directly supervised employees through AI-supported systems. Participants came from fintech, e-commerce, telecommunications, banking and insurance, logistics, and health-technology organizations. Inclusion criteria were at least two years of managerial or specialist experience, direct exposure to AI-supported organizational decision-making, and willingness to participate in an in-depth interview. Participants who had only general familiarity with AI but no practical experience of AI-supported management processes were excluded.

Sampling began purposively and then proceeded theoretically. Initial participants were selected because they had direct experience with AI-supported managerial systems. As coding progressed, theoretical sampling was used to recruit participants who could clarify emerging categories, especially accountability ambiguity, employee resistance, managerial reliance on AI recommendations, and governance mechanisms. Interviews continued until theoretical saturation was achieved. Saturation was reached after 23 interviews, when no substantially new categories emerged, and three additional interviews were conducted to confirm category density and refine the relationships among concepts. The final sample included 15 men and 11 women. Participants ranged in age from 29 to 54 years. In terms of education, 5 participants held bachelor's degrees, 16 held master's degrees, and 5 held doctoral or professional degrees. Regarding organizational role, 13 participants held managerial positions, 7 held HR or people-analytics roles, and 6 held technical, compliance, or AI-governance-related roles.

Data were collected through semi-structured interviews. The interview protocol included open-ended questions about participants' experiences with AI-supported decision-making, changes in managerial authority, the perceived credibility of AI outputs, situations in which managers accepted or rejected algorithmic recommendations, employee reactions to algorithmic decisions, accountability for AI-supported outcomes, and the competencies required for leadership in AI-augmented organizations. Examples of interview questions included: "Can you describe a situation in which an AI system influenced a managerial decision?", "How has AI changed the way authority is exercised in your organization?", "Who is considered responsible when an AI-supported decision creates a problem?", "How do employees respond to algorithmic evaluations or recommendations?", and "What skills do managers need in order to lead effectively with AI systems?" The interview guide was refined during the study as new categories emerged.

The interviews lasted between 45 and 82 minutes, with an average duration of approximately 63 minutes. Interviews were conducted in private meeting rooms or through secure online platforms, depending on participant preference and organizational access. All interviews were audio-recorded with informed consent and transcribed verbatim. Field notes were written after each interview to capture contextual impressions, nonverbal cues, analytical questions, and early conceptual insights. Participants were informed about the purpose of the study, the voluntary nature of participation, confidentiality procedures, and their right to withdraw. To protect confidentiality, participants were identified by codes such as P01, P02, and P03, and organizational names were removed from the data.

Data analysis was conducted using open, axial, and selective coding. In the open-coding stage, interview transcripts were examined line by line to identify initial codes related to managerial decision-making, algorithmic recommendations, trust, control, resistance, accountability, and leadership adaptation. In the axial-coding stage,

conceptually related codes were grouped into categories and subcategories by examining conditions, actions/interactions, and consequences. In the selective-coding stage, the core category was identified and integrated with the main categories to construct the grounded-theory model. Constant comparison was used throughout the analysis to compare incidents with incidents, codes with codes, codes with categories, and categories with emerging theoretical relationships. NVivo software was used to manage transcripts, code segments, retrieve quotations, compare participant groups, and organize analytical memos.

Several strategies were used to enhance trustworthiness. Credibility was supported through prolonged engagement with the data, member checking with selected participants, and comparison of perspectives across managerial, HR, technical, and compliance roles. Dependability was strengthened by maintaining an audit trail that documented sampling decisions, coding revisions, category development, and analytical memos. Confirmability was supported by reflexive memo writing, which helped distinguish participant meanings from researcher assumptions. Transferability was addressed through thick description of the organizational context, participant characteristics, AI applications, and category relationships. The final grounded-theory model was reviewed against the full dataset to ensure that the proposed theory was grounded in the participants' accounts.

3. Findings and Results

The demographic characteristics of the participants showed that the study included 26 individuals working in AI-augmented organizations in Tehran. Of these participants, 15 were men and 11 were women. The age distribution showed that 7 participants were between 29 and 34 years old, 9 were between 35 and 40 years old, 6 were between 41 and 46 years old, and 4 were between 47 and 54 years old. Regarding education, 5 participants held bachelor's degrees, 16 held master's degrees, and 5 held doctoral or professional degrees. In terms of organizational role, 6 participants were senior managers, 7 were middle managers, 5 were HR or people-analytics specialists, 4 were AI or data-product managers, 2 were compliance or risk managers, and 2 were operational team leaders. Participants came from fintech organizations ($n = 6$), e-commerce and platform-based companies ($n = 5$), telecommunications and ICT firms ($n = 5$), banking and insurance organizations ($n = 4$), logistics companies ($n = 3$), and health-technology organizations ($n = 3$). Their work experience ranged from 5 to 24 years, and their direct experience with AI-supported management systems ranged from 1 to 7 years.

The first main category was **algorithmic delegation of managerial judgment**. Participants described how AI systems increasingly performed tasks that had previously been understood as managerial judgment, including prioritizing operational problems, identifying high-risk employees or customers, recommending staffing levels, ranking job applicants, predicting team performance, and detecting deviations from expected productivity patterns. However, participants did not describe this delegation as complete replacement. Rather, AI systems were viewed as producing a first layer of judgment that managers then accepted, modified, delayed, or challenged. A middle manager in a fintech company explained, "The system does not fire anyone, but it tells us who is becoming a performance risk. After that, the manager has to decide whether the number reflects reality or whether there is something the model cannot see" (P07). This category included the subcategories of automated prioritization, predictive recommendation, dashboard-mediated attention, and managerial override. The important finding was that managerial authority shifted from directly observing and evaluating employees to interpreting algorithmically produced signals. Managers increasingly acted after the AI system had already structured the field of attention.

The second main category was **redistribution of authority and accountability ambiguity**. Participants repeatedly emphasized that AI systems changed who appeared to have authority while leaving responsibility unclear. In some organizations, algorithmic recommendations were treated as highly authoritative because they were associated with data, objectivity, and senior-management investment. Yet when a decision produced negative consequences, responsibility returned to human managers. One HR specialist stated, “When the algorithm recommends someone for promotion, everyone says the decision is scientific. But if that person fails later, no one blames the algorithm. They ask why HR and the manager trusted it” (P12). This category included the subcategories of hidden decision ownership, managerial dependence on AI outputs, upward accountability, and responsibility without full control. Participants reported that managers sometimes felt pressured to follow AI recommendations because rejecting them required justification, while accepting them appeared administratively safer. However, this created a paradox: the more authoritative the AI system became, the more ambiguous human accountability felt.

The third main category was **datafied visibility and algorithmic control**. Participants described AI-augmented organizations as environments in which work became more visible, measurable, and comparable. Dashboards, productivity scores, customer-interaction analytics, workflow tracking, and automated alerts created a continuous stream of data about employee behavior and team performance. Some managers viewed this visibility positively because it reduced information asymmetry and enabled earlier intervention. Others believed it intensified surveillance and changed employee behavior in unintended ways. A team leader in an e-commerce organization explained, “Before, I knew my team through conversations and results. Now I also know them through indicators. The problem is that people start working for the indicators, not always for the real customer” (P18). This category included the subcategories of quantified performance, continuous monitoring, behavioral adaptation, metric gaming, and emotional pressure. Participants reported that algorithmic visibility could strengthen managerial coordination, but it could also narrow leadership attention to measurable outputs and weaken sensitivity to informal, relational, and contextual aspects of work.

The fourth main category was **leadership sensemaking and algorithmic literacy**. Participants argued that managers needed new competencies to lead effectively in AI-augmented organizations. Technical expertise alone was not considered sufficient; instead, managers needed enough algorithmic literacy to understand what AI systems could and could not do, how models were trained, what data limitations existed, and when outputs required contextual interpretation. At the same time, they needed communicative competence to explain AI-supported decisions to employees. A senior manager stated, “The leader of the future is not the person who knows coding better than everyone. It is the person who can ask the right question from the system, understand the limitation of the answer, and explain the decision to people” (P03). This category included the subcategories of interpretive competence, critical questioning, translation between technical and managerial language, ethical sensitivity, and employee reassurance. Participants emphasized that employees were more likely to accept AI-supported decisions when managers could explain the logic, acknowledge uncertainty, and show that human judgment remained present.

The fifth main category was **governance of algorithmic legitimacy**. Participants indicated that algorithmic leadership became legitimate only when organizations established procedures for transparency, contestability, human review, and ethical accountability. AI systems were more trusted when employees knew how recommendations were used, when managers could override outputs, when affected employees could appeal

decisions, and when technical teams regularly audited model performance. A compliance manager explained, “The question is not whether the model is accurate only. The question is whether people have a path to challenge it when the output affects their career” (P22). This category included the subcategories of human-in-the-loop governance, explainability routines, audit trails, appeal mechanisms, and participatory system design. Participants viewed governance as essential because AI-supported authority could easily become opaque. Without governance, algorithmic leadership was experienced as imposed control; with governance, it was more likely to be perceived as a disciplined partnership between human judgment and machine intelligence.

The core category that integrated the findings was **calibrated algorithmic authority**. This category refers to the ongoing process through which managers determine the appropriate level of reliance on AI systems while preserving human responsibility, contextual judgment, and organizational legitimacy. Calibrated algorithmic authority emerged through four interconnected processes: delegating selected judgments to AI, interpreting algorithmic outputs through managerial experience, communicating decisions to employees, and governing the boundaries of AI influence. Participants did not reject AI-supported leadership, but they resisted the idea that algorithmic outputs should become unquestionable commands. The grounded theory suggests that algorithmic leadership emerges when managers become calibrators of authority: they decide when AI should inform, when it should recommend, when it should trigger intervention, and when it should be overridden. The outcome of this calibration is not simply better efficiency, but a new form of managerial authority that depends on the balance between data-driven precision and human legitimacy.

4. Discussion and Conclusion

The findings of this study show that algorithmic leadership emerges through the reconstruction of managerial authority in AI-augmented organizations. The core category, calibrated algorithmic authority, indicates that managers do not simply lose authority to AI systems; rather, their authority is transformed. They become responsible for interpreting algorithmic outputs, determining the boundaries of machine influence, justifying decisions to employees, and maintaining legitimacy when decisions are partially produced by nonhuman systems. This finding aligns with Raisch and Krakowski’s (2021) argument that automation and augmentation are interdependent in management. AI may automate selected managerial tasks, but the automation of one task often increases the need for human interpretation, coordination, and ethical judgment in another. Therefore, algorithmic leadership should not be understood as a transition from human leadership to machine leadership, but as the emergence of hybrid authority in which leadership is distributed across human actors, AI systems, data infrastructures, and governance routines.

The first finding, algorithmic delegation of managerial judgment, supports previous research showing that algorithms increasingly perform managerial functions. Parent-Rochelleau and Parker (2022) argue that algorithmic management systems can perform functions such as monitoring, goal setting, performance management, scheduling, compensation, and job termination. The present study extends this literature by showing how managers experience such delegation not merely as task automation but as a reorganization of managerial judgment. In the participants’ accounts, AI systems often produced the first interpretation of organizational reality by identifying risk, ranking alternatives, or flagging deviations. This finding is also consistent with Faraj et al. (2018), who argue that learning algorithms reshape expertise and coordination. In AI-augmented organizations,

expertise is no longer located only in managerial experience; it is partly embedded in predictive systems and dashboards. However, the findings also show that algorithmic judgment remains incomplete because it requires human contextualization. This supports Jarrahi's (2018) view that human–AI symbiosis is especially important in complex and equivocal decision environments.

The second finding, redistribution of authority and accountability ambiguity, contributes to debates about responsibility in AI-supported organizations. Participants described a paradox in which AI recommendations carried strong practical authority, yet human managers remained responsible for consequences. This finding aligns with Tambe et al. (2019), who identify accountability, fairness, and employee reaction as major challenges in AI-based HRM. It also resonates with Shrestha et al. (2019), who argue that organizations must design appropriate decision structures for combining human and AI judgment. The present study adds that decision structures are not merely technical arrangements; they are authority arrangements. When AI outputs are treated as objective but cannot be held accountable, managers may become responsible for decisions they did not fully generate. This creates a legitimacy problem that organizations must address through explicit governance, clear decision rights, and documented override procedures.

The third finding, datafied visibility and algorithmic control, confirms and extends the literature on algorithms as mechanisms of organizational control. Kellogg et al. (2020) conceptualize algorithms at work as a new contested terrain because they reshape how control is enacted, resisted, and negotiated. The current findings show that datafied visibility can improve coordination by making performance patterns more observable, but it can also narrow managerial attention and intensify employee anxiety. This dual effect is important. Managers valued dashboards because they reduced uncertainty and provided early warning signals, yet they also recognized that employees might adapt strategically to metrics, focus on measurable indicators, or experience continuous monitoring as surveillance. This finding supports Lee et al. (2015), who found that algorithmic and data-driven management affects human workers' experiences and work practices. It also aligns with Schildt's (2017) argument that big data and algorithmic systems reshape organizational design through computer-augmented transparency.

The fourth finding, leadership sensemaking and algorithmic literacy, shows that the competencies required for leadership are changing. Participants did not define algorithmic leadership as technical mastery alone. Instead, they emphasized interpretive competence, critical questioning, ethical sensitivity, and communication. This finding aligns with digital leadership literature, which argues that leaders must operate effectively in technologically mediated environments (Avolio et al., 2014; Cortellazzo et al., 2019). However, the present study extends digital leadership theory by showing that AI-augmented leadership requires a specific form of algorithmic literacy. Managers must understand the limits of AI outputs, the assumptions built into data models, and the difference between prediction and judgment. They must also explain AI-supported decisions in ways that employees perceive as meaningful and fair. This finding is consistent with Glikson and Woolley's (2020) review of trust in AI, which shows that trust is shaped by perceived capability, representation, and task characteristics. In organizations, trust in AI is mediated by managerial explanation and sensemaking.

The fifth finding, governance of algorithmic legitimacy, indicates that algorithmic leadership becomes acceptable only when AI-supported authority is transparent, contestable, and accountable. Participants did not view accuracy as sufficient. They emphasized the importance of appeal mechanisms, human review, audit trails, and employee voice. This finding aligns with Binns et al. (2018), who show that perceptions of justice in algorithmic decisions are shaped by concerns about arbitrariness, generalization, and dignity. It also supports

broader arguments that AI governance must address not only technical performance but also social legitimacy. In this study, governance operated as the bridge between algorithmic authority and human acceptance. When employees could not understand or challenge AI-supported decisions, algorithmic leadership appeared coercive. When managers could explain, review, and adjust AI outputs, algorithmic leadership became more legitimate.

The grounded theory developed in this study suggests that algorithmic leadership is best understood as a process of calibration. Calibration means determining the appropriate degree of algorithmic influence in a given managerial situation. In some situations, such as detecting operational anomalies or analyzing large-scale customer patterns, high reliance on AI may be appropriate. In other situations, such as evaluating employee potential, resolving conflict, or making ethically sensitive decisions, AI outputs may need to remain advisory. This distinction reflects Jarrahi's (2018) argument that AI and human intelligence have complementary strengths. It also reflects Haesevoets et al.'s (2021) finding that managers often prefer human-machine partnerships in which humans retain the upper hand. In the present study, participants did not reject AI; they rejected uncalibrated AI authority. Their accounts suggest that effective algorithmic leadership requires managers to know when to trust AI, when to question it, when to explain it, and when to override it.

The findings also have implications for leadership theory. Traditional leadership theories often assume that influence is exercised by human agents. Even distributed leadership approaches tend to distribute leadership among people rather than between people and intelligent systems. The present study suggests that leadership theory must account for nonhuman systems that shape attention, decision-making, control, and legitimacy. AI does not possess leadership in the human sense of intention, morality, or relational awareness, but it can participate in leadership processes by structuring what managers see, what options they consider, what actions are prioritized, and how employees are evaluated. Therefore, algorithmic leadership should be conceptualized as a sociotechnical process in which leadership influence emerges through the interaction of human judgment, algorithmic calculation, organizational rules, and employee interpretation.

The study also contributes to the literature on authority. Managerial authority has historically been grounded in formal position, expertise, organizational hierarchy, and legitimacy. In AI-augmented organizations, authority becomes partly datafied. Decisions gain persuasive force because they appear to be supported by predictive analytics, objective indicators, or machine-generated recommendations. Yet this datafied authority can be fragile because employees may question the fairness, accuracy, or contextual sensitivity of AI systems. Therefore, algorithmic authority requires continuous legitimation. Managers must not only use AI systems; they must actively govern the meaning of AI-supported decisions. This is why the role of the manager becomes more, not less, important in certain respects. The manager becomes the interpreter of machine output, the guardian of contextual judgment, and the communicator of procedural fairness.

The limitations of this study should be acknowledged. First, the study was conducted in Tehran-based organizations, and the findings may reflect the institutional, cultural, technological, and managerial characteristics of that context. Second, the study focused on managers and organizational specialists rather than frontline employees, so the employee experience of algorithmic leadership was captured indirectly. Third, the study examined organizations that had already adopted AI-supported systems, which may have excluded organizations with failed or rejected AI implementation experiences. Fourth, because the study used a qualitative grounded theory design, the findings are intended to generate theory rather than produce statistically generalizable

conclusions. Finally, algorithmic leadership is a rapidly evolving phenomenon, and changes in AI capabilities, regulation, and organizational adoption may alter the dynamics identified in this study.

Future research should examine algorithmic leadership across different industries, countries, and organizational sizes to assess how contextual factors shape the calibration of algorithmic authority. Comparative studies could explore whether algorithmic leadership differs between highly regulated sectors such as banking and healthcare and more flexible sectors such as e-commerce or digital platforms. Future research should also include frontline employees to better understand how algorithmic authority is experienced by those subject to AI-supported evaluation, monitoring, or task allocation. Longitudinal studies would be especially valuable because the relationship between managers and AI systems may change over time as trust develops, failures occur, governance systems mature, and employees learn how to respond to algorithmic metrics. Quantitative research could also develop and validate measurement scales for algorithmic leadership, algorithmic authority calibration, and perceived legitimacy of AI-supported managerial decisions.

In practice, organizations should avoid treating AI implementation as a purely technical project. AI-supported management systems should be introduced with clear decision rights, transparent communication, employee consultation, appeal mechanisms, and human review procedures. Managers should be trained not only in using AI dashboards but also in understanding model limitations, asking critical questions, identifying bias risks, and explaining AI-supported decisions to employees. Organizations should define which decisions can be automated, which require human approval, and which should remain primarily human because of ethical, relational, or contextual sensitivity. Effective algorithmic leadership requires a governance culture in which AI is powerful enough to improve decision quality but not so powerful that it becomes unquestionable. The practical challenge is not to choose between human leadership and AI systems, but to design a disciplined partnership in which algorithmic intelligence strengthens rather than weakens human responsibility.

Ethical Considerations

All procedures performed in this study were under the ethical standards.

Acknowledgments

Authors thank all who helped us through this study.

Conflict of Interest

The authors report no conflict of interest.

Funding

According to the authors, this article has no financial support.

References

- Avolio, B. J., Kahai, S., & Dodge, G. E. (2001). E-leadership: Implications for theory, research, and practice. *The Leadership Quarterly*, 11(4), 615–668. [https://doi.org/10.1016/S1048-9843\(00\)00062-X](https://doi.org/10.1016/S1048-9843(00)00062-X)
- Avolio, B. J., Sosik, J. J., Kahai, S. S., & Baker, B. (2014). E-leadership: Re-examining transformations in leadership source and transmission. *The Leadership Quarterly*, 25(1), 105–131. <https://doi.org/10.1016/j.leaqua.2013.11.003>

- Binns, R., Van Kleek, M., Veale, M., Lyngs, U., Zhao, J., & 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>
- Charmaz, K. (2014). *Constructing grounded theory* (2nd ed.). Sage.
- Corbin, J., & Strauss, A. (2015). *Basics of qualitative research: Techniques and procedures for developing grounded theory* (4th ed.). Sage.
- Cortellazzo, L., Bruni, E., & Zampieri, R. (2019). The role of leadership in a digitalized world: A review. *Frontiers in Psychology*, 10, 1938. <https://doi.org/10.3389/fpsyg.2019.01938>
- Davenport, T. H., & Ronanki, R. (2018). Artificial intelligence for the real world. *Harvard Business Review*, 96(1), 108–116.
- Faraj, S., Pachidi, S., & Sayegh, K. (2018). Working and organizing in the age of the learning algorithm. *Information and Organization*, 28(1), 62–70. <https://doi.org/10.1016/j.infoandorg.2018.02.005>
- Gioia, D. A., Corley, K. G., & Hamilton, A. L. (2013). Seeking qualitative rigor in inductive research: Notes on the Gioia methodology. *Organizational Research Methods*, 16(1), 15–31. <https://doi.org/10.1177/1094428112452151>
- Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of Management Annals*, 14(2), 627–660. <https://doi.org/10.5465/annals.2018.0057>
- Haesevoets, T., De Cremer, D., Dierckx, K., & Van Hiel, A. (2021). Human-machine collaboration in managerial decision making. *Computers in Human Behavior*, 119, 106730. <https://doi.org/10.1016/j.chb.2021.106730>
- Jarrahi, M. H. (2018). Artificial intelligence and the future of work: Human–AI symbiosis in organizational decision making. *Business Horizons*, 61(4), 577–586. <https://doi.org/10.1016/j.bushor.2018.03.007>
- 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>
- Lee, M. K., Kusbit, D., Metsky, E., & Dabbish, L. (2015). Working with machines: The impact of algorithmic and data-driven management on human workers. *Proceedings of the 33rd Annual ACM Conference on Human Factors in Computing Systems*, 1603–1612. <https://doi.org/10.1145/2702123.2702548>
- Parent-Rochelleau, X., & Parker, S. K. (2022). Algorithms as work designers: How algorithmic management influences the design of jobs. *Human Resource Management Review*, 32(3), 100838. <https://doi.org/10.1016/j.hrmr.2021.100838>
- 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>
- Schildt, H. (2017). Big data and organizational design: The brave new world of algorithmic management and computer augmented transparency. *Innovation: Organization & Management*, 19(1), 23–30. <https://doi.org/10.1080/14479338.2016.1252043>
- Shrestha, Y. R., Ben-Menahem, S. M., & von Krogh, G. (2019). Organizational decision-making structures in the age of artificial intelligence. *California Management Review*, 61(4), 66–83. <https://doi.org/10.1177/0008125619862257>
- Tambe, P., Cappelli, P., & Yakubovich, V. (2019). Artificial intelligence in human resources management: Challenges and a path forward. *California Management Review*, 61(4), 15–42. <https://doi.org/10.1177/0008125619867910>
- Zuboff, S. (1988). *In the age of the smart machine: The future of work and power*. Basic Books.