

When Algorithms Become Authority: Digital Silence in AI-Mediated HRM

Nancy Vishwas¹, Sheeba Asirvatham²

^{1,2}Department of Management Studies, Woosong University, Daejeon, Republic of Korea

Abstract

As a modern Organizations embed artificial intelligence (AI) into human resource management (HRM) functions such as recruitment, performance evaluation, and workforce analytics. A distinct behavioural challenge emerges. Employees frequently observe flawed, or AI driven algorithmic output but choose to remain silent. This article conceptualises digital silence, defined as the deliberate withholding of concerns, feedbacks, observation regarding Ai driven decision. Drawing upon organizational silence theory, the Ability-Motivation-Opportunity (AMO) framework and algorithmic management theory we develop a comprehensive conceptual model illustrating how supportive human resource practices (HRP) alleviate digital silence to cultivate sustained institutional trust in AI system. We theorise algorithm opacity as critical boundary condition and explore how remote work voice architectures shape these dynamic. Four explicit theoretical proposition are advanced to guide future empirical research.

1. Introduction

The integration of artificial intelligence (AI) into routine human resource management practices has evolved from an emerging technical trend into a central operational strategy. Algorithmic system now screen resumes, evaluate employee productivity and inform promotion decisions under the premise of optimising efficiency and mitigating human bias. However, these systems often reduce qualitative human experience to narrow quantitative metrics a phenomenon known as algorithmic reductionism which frequently damages perceived procedural justice. While existing literature examines algorithm aversion and fairness perceptions it overlooks a critical behavioural response: why employees stay silent when they notice an AI system making an unfair decision.

The organisational silence assume human to human workplace hierarchies where withholding voice is predominantly motivated by fear of supervisor retaliation, hierarchical distance or relational damage. Challenging an AI system involves a fundamentally different psychological mechanism. Facing statistical models, employees experience self doubt, questioning whether they possess the technical authority or standing to challenge a system perceived as objective and authoritative. This article addresses this gap by defining digital silence as a distinct behavioural construct and providing a theoretical model to explain its organisational drivers, boundary conditions and downstream impacts.

2. Conceptualizing Digital Silence as a Distinct Organizational Behavior Construct

2.1 Definition

Organizational silence, as conceptualized by Morrison and Milliken (2000), describes employees' collective tendency to withhold opinions and concerns about organizational problems, shaped by perceived psychological safety and institutional signals. Van Dyne, Ang, and Botero (2003) subsequently

refined this view, showing that silence is not simply the behavioral opposite of voice but a multidimensional construct in its own right, driven by acquiescent, defensive, or prosocial motives. This article defines Digital Silence as the intentional withholding of concerns, observations, or feedback regarding AI-driven decisions in the workplace.

2.2 Why Digital Silence Is a Distinct Construct

A reasonable objection to this article's central claim is that Digital Silence is simply organizational silence occurring in a technologically mediated setting – the same phenomenon with a new label attached. This subsection addresses that objection directly by examining four dimensions on which Digital Silence departs from organizational silence and its closest relatives: the target of the withheld communication, the underlying psychology that produces the withholding, the downstream consequences of sustained silence, and the HRM implications that follow.

Different target

Classical organizational silence, employee voice, and whistleblowing theories explicitly assume a human recipient capable of empathy, negotiation, and relational repair (Morrison & Milliken, 2000; Van Dyne et al., 2003). In contrast, Digital Silence targets automated systems where relational feedback is displaced onto diffuse, often inaccessible human reviewers or third-party vendor networks (Duggan et al., 2026). Decoupling the decision-making source from the immediate supervisory line creates a structural novelty of algorithmic management that traditional hierarchical silence frameworks fail to capture (Kellogg et al., 2020).

Different psychology

The psychological driver of classical organizational silence is predominantly interpersonal—fear of damaging a relationship or inviting supervisor retaliation (Morrison & Milliken, 2000). The psychological driver of Digital Silence is substantially epistemic. Facing an opaque scoring model, employees do not necessarily fear interpersonal conflict; rather, they experience epistemic self-doubt, questioning whether they possess the technical vocabulary, validity, or standing to challenge a system perceived as objective and statistically grounded (Bartosiak & Modlinski, 2022; Dietvorst et al., 2015).

Different consequences

While traditional organizational silence resolution depends on shifting human dynamics within a hierarchy, an AI system continuously executes its programmed logic irrespective of unexpressed employee doubts. Sustained Digital Silence severs the operational feedback loops required for organizational error correction (Lee & See, 2004). Unlike whistleblowing, which typically involves episodic reporting of discrete misconduct, Digital Silence represents a diffuse, cumulative accumulation of unvoiced concerns that systematically erodes automation trust over time.

Different HR implications

Finally, psychological safety initiatives reduce interpersonal fear but leave epistemic barriers unresolved; an employee may feel entirely safe to speak yet remain silent due to a lack of technical confidence. Applying the Ability-Motivation-Opportunity (AMO) framework (Appelbaum et al., 2000), addressing Digital Silence mandates building Ability (AI literacy) alongside Motivation (psychological safety) and Opportunity (low-risk voice channels).

Table 1. Digital Silence Relative to Adjacent Constructs

Construct	Object of the behavior	Underlying motive	Directionality	Typical outcome studied
Organizational silence (Morrison & Milliken, 2000)	Human supervisors, management, organizational policy	Fear of interpersonal or hierarchical repercussions	Withholding directed at people with formal authority	Organizational change, climate
Employee voice/silence (Van Dyne, Ang, & Botero, 2003)	Any work-related issue or improvement opportunity	Acquiescent, defensive, or prosocial motives	General expression vs. withholding of ideas	Performance, engagement
Whistleblowing (Near & Miceli tradition)	Perceived wrongdoing or misconduct	Moral obligation to stop a harmful practice	Reporting to internal or external authorities	Legal/ethical compliance
Algorithm aversion (Dietvorst, Simmons, & Massey, 2015)	An algorithm's output or forecast	Loss of confidence after witnessing algorithmic error	Reduced reliance / abandonment of algorithm use	Adoption, reliance behavior
Digital Silence (this article)	AI-driven organizational decisions	Epistemic self-doubt and perceived illegitimacy of challenging an opaque system	Withholding of concerns specifically about algorithmic outputs, while continuing to use the system	Trust in AI, algorithmic feedback loops

3. Theoretical Foundations

3.1 Organizational Silence Theory

Morrison and Milliken's (2000) organizational silence theory establishes that silence is an active, context-dependent response shaped by organizational climate and perceived safety. Van Dyne et al. (2003) expanded this framework by identifying acquiescent, defensive, and prosocial motives underlying silence. This article extends these principles to AI-driven environments, conceptualizing Digital Silence as a blend of acquiescent motives (deference to perceived algorithmic authority) and defensive motives (fear of appearing technically unqualified).

3.2 The Ability-Motivation-Opportunity (AMO) Framework

The AMO framework (Appelbaum, Bailey, Berg, & Kalleberg, 2000) offers a useful mechanism for mapping how specific HRM practices are theorized to reduce Digital Silence. AI literacy training is conceptualized as building employees' Ability to recognize and articulate algorithmic problems without which an employee may sense that something is wrong but lack the vocabulary or technical grounding to

raise it credibly. Psychological safety (Edmondson, 1999) and transparent, non-punitive appraisal practices are conceptualized as building Motivation to speak up despite uncertainty about whether one's concern is valid. Participative HR policies escalation channels structured feedback mechanisms, cross-functional review forums are conceptualized as creating the Opportunity to voice concerns about AI outputs through a defined, low-risk channel. Together, these three pathways are proposed to jointly counteract the barriers that sustain Digital Silence; the AMO framework's contribution here is to show that HRM practices are not fungible substitutes for one another but map onto genuinely distinct psychological barriers.

3.3 Algorithmic Management Theory

Algorithmic management theory explains how data-driven systems reconfigure authority, surveillance, and control in the modern workplace (Kellogg, Valentine, & Christin, 2020). This literature clarifies why Digital Silence emerges in the first place: algorithmic authority is diffuse, often opaque, and frequently perceived as beyond employee influence – precisely the condition under which traditional voice mechanisms, designed for human hierarchies, may prove insufficient without deliberate adaptation. Meijerink, Boons, Keegan, and Marler (2021) synthesize this fast-growing literature under the umbrella term “algorithmic HRM,” defined as the use of software algorithms operating on digital data to augment or automate HR decisions, and note that such systems are diffusing rapidly across recruitment, performance management, and workforce analytics alike. Recent reviews of algorithmic management continue to document this diffusion and call for renewed attention to worker-level behavioral responses rather than system-level performance alone (Zhang, Cooke, Ahlstrom, & McNeil, 2025). Research on algorithmic reductionism further shows that employees and applicants often perceive AI-driven HR decisions as stripping away contextual nuance in favor of narrow, quantifiable inputs, which can simultaneously lower perceived fairness and raise employees' uncertainty about whether a qualitative objection would even be heard (Newman, Fast, & Harmon, 2020), a pattern consistent with workers' broader tendency to judge AI-driven HR decisions as less dignity-respecting than human ones (Bankins, Formosa, Griep, & Richards, 2022). In platform-based work specifically, algorithmic management has been shown to actively silence worker voice by removing the human intermediaries who would otherwise hear a concern (Duggan, Dasgupta, McDonnell, Carbery, & Sherman, 2026).

3.4 Digital Voice Architectures and the Changing Context of Voice

A further contextual shift strengthens the case for treating Digital Silence as distinct from classical organizational silence: the physical and relational conditions under which voice historically occurred have themselves been transformed by remote and hybrid work. Kalfa, Kwon, Prouska, and Wilkinson (2025) argue that employee voice research has traditionally assumed co-located, synchronous work environments, and that psychological, temporal, structural, and technological distance in remote and hybrid arrangements creates new barriers to voice that digital channels do not automatically resolve. Their account suggests that digital communication platforms, while nominally providing a channel for feedback, can paradoxically make speaking up feel more formal, more permanent, and more easily monitored than an informal hallway conversation once was raising the perceived risk of voice even when a channel technically exists. This article incorporates that insight as a structural amplifier of Digital Silence: an employee who might otherwise report a concern about an AI-driven decision may be further discouraged by a digital feedback architecture that renders every report visible, timestamped, and attributable, compounding the epistemic hesitation described in Section 2.2 with a distinctly structural one. This

underscores that the Opportunity pathway in the AMO framework (Section 3.2) is not simply a matter of whether a channel exists, but of how psychologically safe that channel feels to use.

3.5 Trust: From Interpersonal Trust to Trust in Automation

Because trust in AI is the outcome variable in this model, it warrants its own theoretical grounding. Mayer, Davis, and Schoorman's (1995) integrative model of organizational trust conceptualizes trust as a willingness to be vulnerable to another party based on perceptions of that party's ability, benevolence, and integrity. Extending trust theory to non-human agents, Lee and See (2004) argue that trust in automation is shaped by the same underlying psychological processes as interpersonal trust, but is calibrated through direct experience with a system's performance, reliability, and predictability over time. Glikson and Woolley's (2020) review of empirical research on human trust in AI similarly concludes that trust is shaped jointly by cognitive factors, such as perceived reliability and transparency, and affective factors, such as comfort and anthropomorphism, and that these determinants interact differently depending on the form and embodiment of the AI system. This article draws on these traditions to argue that Digital Silence disrupts the calibration process Lee and See describe: when employees do not report perceived errors, the organization loses the corrective feedback that would normally allow trust in the AI system to be recalibrated toward an accurate, evidence-based level, leaving employees' trust to drift downward privately rather than being resolved through visible correction.

3.6 Contemporary Developments: Explainability, Governance, and AI Trust (2021–2026)

Recent scholarship has begun to converge on transparency and explainability as central levers for building trust in workplace AI, offering further theoretical support for this article's emphasis on algorithmic opacity as a boundary condition. In a qualitative study of trust-building requirements for successful AI adoption, Bedüe and Fritzsche (2022) find that organizations which fail to address employees' trust requirements directly experience persistently slow and incomplete AI adoption, regardless of a system's technical performance. In a large-scale field study, Yang, Lu, and Cooke (2026) show that explainable-AI interventions particularly counterfactual and locally-scoped explanations measurably increase workers' acceptance of algorithmic decisions and improve perceived management relations, though they also find that overly complex combinations of explanation types can overwhelm workers and erode these gains. This finding is directly relevant to Proposition 4 below: it suggests that reducing algorithmic opacity is not simply a matter of providing more explanation, but of calibrating explanation to what an employee can actually process, a nuance that HRM-side interventions such as AI literacy training are theorized to support. Complementing this, Fenwick, Molnar, and Frangos (2024) argue that HR's evolving role in AI-integrated organizations increasingly centers on “humanizing” AI systems building tools and processes that enhance trust and minimize fear which positions HR not merely as a passive recipient of AI tools but as an active shaper of the conditions under which Digital Silence is more or less likely to occur.

At the organizational level, Papagiannidis, Mikalef, and Conboy (2025) synthesize the responsible-AI-governance literature into a framework built on structural, relational, and procedural governance practices, arguing that responsible AI deployment depends as much on organizational process as on system design. This governance-level perspective complements the individual-level focus of this article: where Papagiannidis et al. (2025) theorize how organizations structure accountability for AI systems, this article theorizes how individual employees behaviorally respond when that accountability structure fails to surface a problem they have noticed. Together, this body of contemporary work reinforces the claim that algorithmic opacity is not a static technical property but a jointly determined outcome of system design, explanation strategy, and organizational governance all of which sit outside the direct control of HRM

practices alone, consistent with this article's positioning of opacity as a boundary condition on HRM effectiveness rather than a variable HRM can fully resolve.

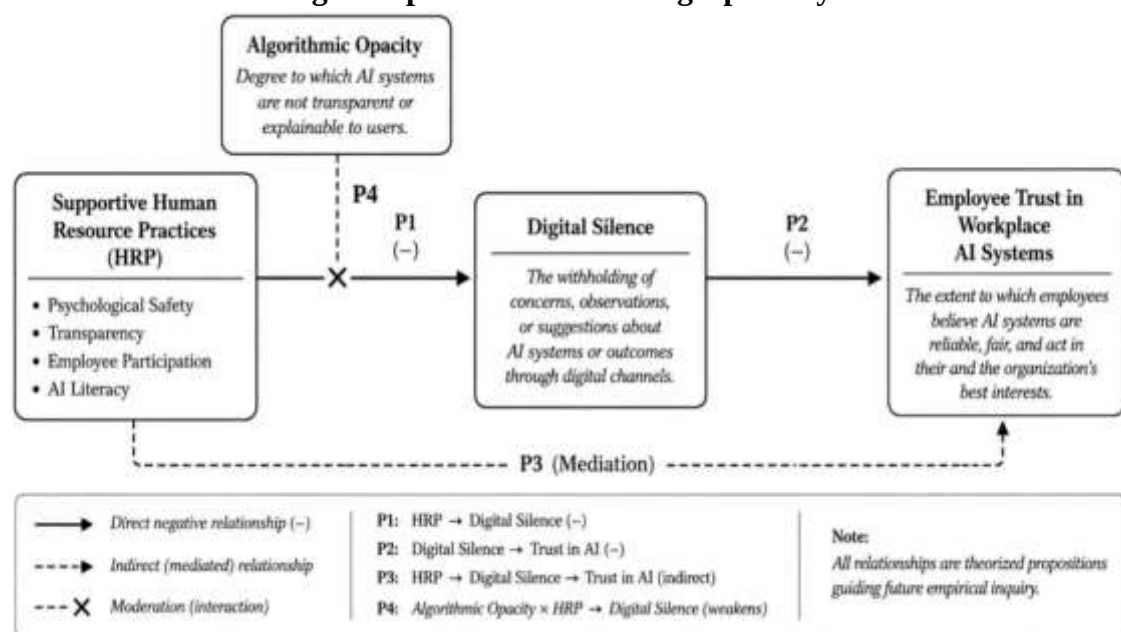
4. Digital Silence and Organizational Outcomes

When Digital Silence is sustained across an organization, its consequences for HRM outcomes are theorized to be substantial. First, employee performance may deteriorate as workers adapt their behavior to conform to AI-driven expectations they privately believe are flawed, reducing discretionary effort and task quality a pattern consistent with experimental evidence that employees often conform to biased algorithmic recommendations rather than actively opposing them (Bartosiak & Modlinski, 2022). Second, organizational learning is theorized to be systematically impaired: when employees withhold observations of system errors, the feedback loops through which organizations identify and correct algorithmic failures are broken, echoing Edmondson's (1999) argument that psychological safety is a precondition for the kind of error reporting that organizational learning depends on. Third, institutional trust in AI is theorized to erode over time as employees accumulate private evidence of AI failures they have never reported, increasing cynicism about AI systems and reducing the behavioral engagement needed to make human-AI collaboration work. This third mechanism is consistent with Lee and See's (2004) view that trust calibration depends on visible, resolved feedback – in its absence, trust does not stabilize at an accurate level but drifts, unobserved by the organization, toward quiet disengagement.

5. Conceptual Model

Building on these theoretical traditions, this article proposes an integrative conceptual model in which supportive HRM practices psychological safety, transparency, participation, and AI literacy are theorized to reduce Digital Silence, which in turn is theorized to shape employees' trust in AI systems. Algorithmic opacity is positioned as a boundary condition that constrains how far HRM practices can suppress silence, since even highly supportive climates may not fully overcome employees' epistemic self-doubt when the system itself is incomprehensible.

Figure 1 presents this model graphically.



Note: Algorithmic Opacity is theorized as a moderator of the HRM Practices → Digital Silence path (P4). Digital Silence is theorized as a mediator of the indirect HRM Practices → Trust in AI relationship (P3).

6. Theoretical Propositions

Because this is a conceptual article, the relationships below are advanced as theoretical propositions intended to guide future empirical inquiry, rather than as hypotheses tested against data. Each proposition is accompanied by the mechanism through which the relationship is theorized to operate:

Proposition 1 (P1): Supportive human resource practices (psychological safety, transparency, employee participation, and AI literacy) are negatively associated with Digital Silence. By addressing ability, motivation, and opportunity simultaneously, organizations systematically diminish the psychological and structural barriers to voice.

Proposition 2 (P2): Digital Silence is negatively associated with overall employee trust in workplace AI systems. Suppressing observations disrupts the feedback loop necessary to correct system errors, replacing active error resolution with compounding private doubts.

Proposition 3 (P3): Digital Silence mediates the relationship between supportive human resource practices and employee trust in AI. Human resource initiatives do not build trust in automated systems directly; rather, they shape trust indirectly by establishing the behavioral path through which algorithmic errors are surfaced and resolved.

Proposition 4 (P4): Algorithmic opacity moderates the negative relationship between human resource practices and Digital Silence, such that the relationship is weaker when opacity is high. High opacity creates a technical ceiling that organizational climate alone cannot overcome, demonstrating that internal human resource policies and system-level explainability investments must function as complementary controls.

7. Discussion

7.1 Implications for Organizational Silence Theory

This article's central theoretical contribution is the introduction of Digital Silence as a construct that extends organizational silence theory into algorithmic workplace contexts while remaining conceptually distinct from adjacent constructs, as established in Section 2.2 and Table 1. For organizational silence theory specifically, this extension implies that the theory's core assumption – that silence is calibrated primarily against interpersonal risk – needs to be supplemented with an epistemic dimension when the object of silence is non-human. Future refinements of silence theory may need to treat “confidence in one's standing to challenge the source” as a distinct antecedent alongside the classical antecedents of psychological safety and perceived futility that Morrison and Milliken (2000) identify, and to account for how the changing architecture of workplace voice itself – increasingly digital, increasingly monitorable – reshapes the perceived risk of speaking up (Kalfa, Kwon, Prouska, & Wilkinson, 2025).

7.2 Implications for HRM Theory and Practice

For HRM theory, this article extends the AMO framework by illustrating how each of its three pathways – ability, motivation, opportunity – maps onto a distinct organizational lever for counteracting a specifically algorithmic form of silence. This reframes AI literacy not as a peripheral digital-skills initiative but as a core voice-enabling HR practice on the same theoretical footing as psychological safety. For HRM practice, the model implies that organizations investing heavily in one AMO pathway while

neglecting the others – for instance, rolling out AI ethics training (Ability) without creating a channel to act on it (Opportunity) – are likely to see only partial reductions in Digital Silence.

7.3 Implications for AI Governance

This article also speaks to the growing responsible-AI-governance literature (Papagiannidis, Mikalef, & Conboy, 2025). That literature has largely theorized governance as a structural and procedural matter – who reviews AI decisions, what documentation is kept, what oversight committees exist. Digital Silence identifies a governance blind spot: even well-designed structural governance can fail if the employees closest to an AI system's day-to-day outputs do not report what they observe. This suggests that responsible AI governance frameworks should explicitly incorporate employee-voice mechanisms – not merely as an HR concern, but as a governance control, since unreported errors are errors a governance structure cannot act on.

7.4 Why This Matters for Future AI-Enabled Workplaces

As AI systems take on a growing share of consequential workplace decisions, and as explainable-AI techniques become more available but not necessarily easier to use well (Yang, Lu, & Cooke, 2026), the gap between what an AI system does and what employees are willing to say about it is likely to become a more, not less, consequential organizational blind spot. Organizations that treat AI trust as purely a function of system accuracy risk missing the behavioral layer this article foregrounds: trust is not simply a reaction to how well a system performs, but a reaction to whether the organization ever finds out when it does not.

8. Practical Implications for Stakeholders

Managerial and HR Practice:

HR practitioners and line managers must deploy balanced AMO practices. Combining AI literacy (Ability) with psychological safety (Motivation) and non-punitive escalation channels (Opportunity) ensures employees possess both the technical grounding and confidence to flag algorithmic errors. Manager-facing training and clear escalation protocols are essential to prevent supervisors from reproducing Digital Silence at higher administrative levels.

System Design and Policy Governance:

AI developers, vendors, and policymakers share responsibility for reducing system opacity. Designers must provide clear, locally-scoped decision rationales to enable informed employee feedback (Yang et al., 2026). Regulators should look beyond static compliance audits by requiring organizations to maintain active internal feedback loops for ongoing system oversight.

9. Directions for Future Empirical Research

1. Develop and validate a multi-item scale for Digital Silence, building on an illustrative item such as “I have noticed issues in AI-driven decisions but chose not to report them,” and establish its discriminant validity relative to existing organizational silence, employee voice, and algorithm aversion measures using confirmatory factor analysis.
2. Test Propositions 1 through 4 using cross-sectional survey data and structural equation modeling, measuring supportive HRM practices as a composite construct, Digital Silence via the newly developed scale, trust in AI via established technology-trust measures adapted from Lee and See's (2004) framework, and algorithmic opacity via established black-box perception measures.

3. Employ longitudinal designs to track how sustained Digital Silence reshapes organizational culture and trust in AI over time, since a single cross-sectional study cannot establish whether HRM practices precede reductions in Digital Silence or whether reverse or reciprocal relationships are also present – for example, whether early instances of unresolved Digital Silence themselves erode employees' perception of how psychologically safe the organization actually is.
4. Examine industry and AI-maturity as boundary conditions, given that the salience of algorithmic opacity and the availability of AI literacy resources likely vary substantially across sectors, and given that algorithm aversion research suggests reactions to algorithmic error may differ by task type (Dietvorst, Simmons, & Massey, 2015).
5. Investigate whether the intensity of Digital Silence varies with the degree of human oversight retained in a decision, comparing AI-assisted, human-in-the-loop, and fully delegated algorithmic decision environments; on theoretical grounds this article expects silence to increase as human intervention is progressively removed, though the functional form of that relationship remains an open empirical question.
6. Investigate managerial and system level interventions (e.g., explainable-AI interfaces, structured escalation protocols, human-in-the-loop review, and the digital voice architectures discussed in Section 3.4) as complements to HRM practices in reducing Digital Silence under conditions of high opacity, building on experimental evidence that explanation type and complexity shape worker acceptance (Yang, Lu, & Cooke, 2026).
7. Explore whether Digital Silence exhibits the same acquiescent/defensive/prosocial subdimensions that Van Dyne, Ang, and Botero (2003) identified for general organizational silence, or whether algorithmic contexts generate a qualitatively different motivational structure.
8. Extend the model to the organizational governance level by examining whether structural, relational, and procedural AI governance practices (Papagiannidis, Mikalef, & Conboy, 2025) moderate the HRM Digital Silence relationship independently of algorithmic opacity.

10. Conclusion

As AI systems assume a greater decision making authority in HRM understanding employee voice behaviour becomes critical to successful technology implementation. This article has introduced digital silence as distinct theoretical construct differentiated it from traditional voice and silence literature across four core dimensions and articulated an integrative model linking supportive HRM practices to automation trust through the mediating role of digital silence and the moderating influence of algorithmic opacity. By addressing both technical and psychological barriers to voice organisation can ensure that human AI collaboration remains transparent, reliable and grounded in mutual trust.

References

1. Appelbaum, E., Bailey, T., Berg, P., & Kalleberg, A. L. (2000). *Manufacturing advantage: Why high-performance work systems pay off*. Cornell University Press.
2. Bankins, S., Formosa, P., Griep, Y., & Richards, D. (2022). AI decision making with dignity? Contrasting workers' justice perceptions of human and AI decision making in a human resource management context. *Information Systems Frontiers*, 24(3), 857–875. <https://doi.org/10.1007/s10796-021-10223-8>

3. Bedüé, P., & Fritzsche, A. (2022). Can we trust AI? An empirical investigation of trust requirements and guide to successful AI adoption. *Journal of Enterprise Information Management*, 35(2), 530–549. <https://doi.org/10.1108/JEIM-06-2020-0233>
4. Bowen, D. E., & Ostroff, C. (2004). Understanding HRM–firm performance linkages: The role of the “strength” of the HRM system. *Academy of Management Review*, 29(2), 203–221. <https://doi.org/10.5465/amr.2004.12736076>
5. Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144(1), 114–126. <https://doi.org/10.1037/xge0000033>
6. Duggan, J., Dasgupta, P., McDonnell, A., Carbery, R., & Sherman, U. (2026). Tensions in algorithmic HRM: Worker voice and organizational silencing in app-based platform work. *International Journal of Human Resource Management*. Advance online publication. <https://doi.org/10.1080/09585192.2026.2699267>
7. Duggan, J., Sherman, U., Carbery, R., & McDonnell, A. (2020). Algorithmic management and app-work in the gig economy: A research agenda for employment relations and HRM. *Human Resource Management Journal*, 30(1), 114–132. <https://doi.org/10.1111/1748-8583.12258>
8. Edmondson, A. (1999). Psychological safety and learning behavior in work teams. *Administrative Science Quarterly*, 44(2), 350–383. <https://doi.org/10.2307/2666999>
9. Fenwick, A., Molnar, G., & Frangos, P. (2024). Revisiting the role of HR in the age of AI: Bringing humans and machines closer together in the workplace. *Frontiers in Artificial Intelligence*, 6, Article 1272823. <https://doi.org/10.3389/frai.2023.1272823>
10. 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>
11. Kalfa, S., Kwon, B., Prouska, R., & Wilkinson, A. (2025). Disconnected workers: Can digital voice fill the gap? *Human Resource Management Journal*, 36(1), 32–47. <https://doi.org/10.1111/1748-8583.70006>
12. Kellogg, K. C., Valentine, M. A., & Christin, A. (2020). Algorithms at work: The new contested terrain of control. *Academy of Management Annals*, 14(1), 366–410. <https://doi.org/10.5465/annals.2018.0174>
13. Kim, S. (2025). Strategic human resource management in the era of algorithmic technologies: Key insights and future research agenda. *Human Resource Management*. Advance online publication. <https://doi.org/10.1002/hrm.22268>
14. 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>
15. LePine, J. A., & Van Dyne, L. (1998). Helping and voice extra-role behaviors: Evidence of construct and predictive validity. *Academy of Management Journal*, 41(1), 108–119. <https://doi.org/10.2307/256902>
16. Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *Academy of Management Review*, 20(3), 709–734. <https://doi.org/10.2307/258792>
17. Meijerink, J., Boons, M., Keegan, A., & Marler, J. (2021). Algorithmic human resource management: Synthesizing developments and cross-disciplinary insights on digital HRM. *International Journal of Human Resource Management*, 32(12), 2545–2562. <https://doi.org/10.1080/09585192.2021.1925326>

18. Meijerink, J., & Bondarouk, T. (2023). The duality of algorithmic management: Toward a research agenda on HRM algorithms, autonomy and value creation. *Human Resource Management Review*, 33(1), Article 100876. <https://doi.org/10.1016/j.hrmr.2021.100876>
19. Melo, N., & Reis, J. (2026). Generative artificial intelligence in human resource management practice. *Administrative Sciences*, 16(3), Article 113. <https://doi.org/10.3390/admsci16030113>
20. Morrison, E. W. (2014). Employee voice and silence. *Annual Review of Organizational Psychology and Organizational Behavior*, 1(1), 173–197. <https://doi.org/10.1146/annurev-orgpsych-031413-091328>
21. Morrison, E. W., & Milliken, F. J. (2000). Organizational silence: A barrier to change and development in a pluralistic world. *Academy of Management Review*, 25(4), 706–725. <https://doi.org/10.2307/259200>
22. Newman, D. T., Fast, N. J., & Harmon, D. J. (2020). When eliminating bias isn't fair: Algorithmic reductionism and procedural justice in human resource decisions. *Organizational Behavior and Human Decision Processes*, 160, 149–167. <https://doi.org/10.1016/j.obhdp.2020.03.008>
23. Papagiannidis, E., Mikalef, P., & Conboy, K. (2025). Responsible artificial intelligence governance: A review and research framework. *Journal of Strategic Information Systems*, 34(2), Article 101885. <https://doi.org/10.1016/j.jsis.2024.101885>
24. Pasquale, F. (2015). *The black box society: The secret algorithms that control money and information*. Harvard University Press.
25. Tinguely, P. N., Lee, J., & He, V. F. (2023). Designing human resource management systems in the age of artificial intelligence. *Management Revue*. Advance online publication. <https://doi.org/10.1007/s41469-023-00153-x>
26. Van Dyne, L., Ang, S., & Botero, I. C. (2003). Conceptualizing employee silence and employee voice as multidimensional constructs. *Journal of Management Studies*, 40(6), 1359–1392. <https://doi.org/10.1111/1467-6486.00384>
27. Yang, M. M., Lu, Y., & Cooke, F. L. (2026). Demystifying AI for the workforce: The role of explainable AI in worker acceptance and management relations. *Journal of Management Studies*, 63(2), 438–472. <https://doi.org/10.1111/joms.70039>
28. Zhang, M. M., Cooke, F. L., Ahlstrom, D., & McNeil, N. (2025). The rise of algorithmic management and implications for work and organisations. *New Technology, Work and Employment*, 40(3), 659–671. <https://doi.org/10.1111/ntwe.12343>