

AI ETHICS IN HUMAN RESOURCE MANAGEMENT: A SYSTEMATIC REVIEW

Dismus ONDERI

*KCA University, School of Business, Nairobi, Kenya
dismusonderi15@gmail.com*

Dorothy KEMUNTO

*KCA University, School of Business, Nairobi, Kenya
dorothykemunto254@gmail.com*

Edwin WANJOHI NYINGI

*KCA University, School of Business, Nairobi, Kenya
eddiewnyingi@gmail.com*

Elizabeth KANAKE

*KCA University, School of Business, Nairobi, Kenya
lizkanake@yahoo.com*

Jackson NDOLO

*KCA University, School of Business, Nairobi, Kenya
jndolo@kcau.ac.ke*

Caroline NTARA

*KCA University, School of Business, Nairobi, Kenya
c.ntara@kcau.ac.ke*

Abstract

Artificial intelligence (AI) is increasingly reshaping key human resource management functions, including recruitment and selection. However, its rapid adoption has often outpaced the development of ethical and regulatory frameworks. This has raised significant concerns, including data privacy risks, algorithmic bias, and potential impacts on employee dignity and limited transparency. This systematic review synthesises empirical and theoretical findings from 12 peer-reviewed articles published between 2019 and 2025. The review employed the PEO (Population, Exposure, Outcome) framework and followed the PRISMA 2020 guidelines for reporting. It identifies five key themes: (1) bias and discrimination in algorithm-based recruitment and selection; (2) transparency, explainability, and accountability; (3) trust, fairness perceptions, and dignity among employees; (4) organisational governance and ethical AI systems; and (5) strategies for reducing bias and promoting inclusive design. The review highlights key ethical risks in AI-driven Human Resource applications and identifies ongoing governance challenges, particularly regarding context-specific implementation in non-Western and African organisational environments. Although the review adopts a broad HRM perspective, the available empirical evidence is predominantly concentrated on recruitment and selection, with comparatively limited research examining ethical issues in other HR functions. It underscores the importance of ethical governance, transparency, and inclusive design in AI-driven HR systems, while highlighting implications for HR practitioners, leaders, policymakers, and researchers. The study also emphasises the need for more context-sensitive regulatory approaches and calls for further empirical research grounded in diverse cultural and institutional settings.

Keywords: artificial intelligence, human resource management, HR accountability, HR ethics

1. INTRODUCTION

The integration of AI into human resource management (HRM)¹ represents a profound reconfiguration of how organisations attract, evaluate, manage, and retain talent. Traditionally grounded in human judgment and relational processes, HRM is increasingly mediated by algorithmic systems that process large-scale data, identify complex patterns, and generate predictive insights. AI applications now permeate nearly every human resource (HR) function, including automated résumé screening, natural language processing in candidate evaluation, workforce analytics, and algorithmic performance management. This diffusion reflects not merely technological adoption but a structural transformation in organisational decision-making. As organisations confront increasingly complex and globalised labour markets, AI-enabled HR tools are becoming integral to achieving efficiency, scalability, and strategic responsiveness (Hunkenschroer and Luetge 2022; Tambe et al. 2019; Zhu 2025).

At its core, the appeal of AI in HRM lies in its promise to enhance efficiency and objectivity. Algorithmic systems can significantly reduce time-to-hire, standardise evaluation criteria, and improve predictive accuracy in identifying high-performing or at-risk employees. In principle, these systems offer a way to minimise human bias by relying on structured data rather than subjective impressions. However, this promise is inherently contingent on the quality, representativeness, and governance of the data used to train such systems. A growing body of research demonstrates that AI systems can reproduce and even amplify existing societal biases embedded within historical datasets, thereby institutionalising discrimination under the guise of technological neutrality (Panarese et al. 2025; Raghavan et al. 2020; Rieke and Bogen 2018; Zhu 2025).

High-profile failures have brought these risks into sharp focus. A widely cited case involves an experimental recruitment system developed by Amazon that systematically downgraded résumés containing indicators of female identity because of biased training data (Dastin 2018). Similarly, controversies surrounding algorithmic evaluation systems have demonstrated how opaque models can produce erroneous and unjust outcomes with significant professional consequences (Cambridge Consultants 2025; Panarese et al. 2025). These cases highlight a critical insight: AI systems are not neutral arbiters, but socio-technical constructs shaped by data, institutional contexts, and design choices. Without deliberate ethical safeguards, such systems risk perpetuating and legitimising discrimination (Hunkenschroer and Luetge 2022; Selbst et al. 2019).

Beyond issues of bias and fairness, the expansion of AI in HRM raises critical concerns related to employee surveillance, autonomy, and workplace experience (Bastida et al. 2025; Rabenu and Baruch 2025). The increasing deployment of digital monitoring tools, including productivity tracking systems, behavioural analytics, and sentiment analysis, has transformed managerial oversight. While these technologies are often justified as mechanisms for enhancing performance and accountability, they also introduce new forms of control that may erode trust (Newman et al. 2020) and employee agency. Emerging research indicates that algorithmic management can contribute to heightened psychological strain, reduced perceived autonomy, and increased job insecurity, particularly when decision-making processes are opaque or lack procedural transparency (Kellogg et al. 2020; Pliszka 2025). These concerns are especially salient in hybrid and digitally mediated work environments, where algorithmic oversight is both pervasive and less visible.

The ethical implications of AI in HRM extend beyond individual organisations to encompass broader questions of governance and accountability. Policymakers and international organisations are increasingly responding to these challenges by developing AI governance frameworks that emphasise transparency, fairness, and human oversight. For example, regulatory initiatives within the European Union and global principles advanced by institutions such as the OECD and UNESCO underscore the importance of human-centred AI (European Commission, 2019; OECD, 2019; UNESCO, 2021). Nevertheless, the pace of technological innovation continues to outstrip regulatory development, leaving organisations to navigate complex ethical terrains with limited standardised guidance (Floridi et al., 2018; Jobin et al., 2019).

¹ This review adopts a broad human resource management (HRM) perspective. However, the current empirical literature on AI ethics in HRM is disproportionately concentrated on recruitment and selection. Accordingly, the synthesis reflects the distribution of the available evidence rather than an intentional emphasis on a particular HR function. The limited empirical research on AI ethics in areas such as performance management, learning and development, compensation, employee monitoring, workforce planning and career development highlights an important gap in the literature and a priority for future research.

AI ETHICS IN HUMAN RESOURCE MANAGEMENT: A SYSTEMATIC REVIEW

In this context, the intersection of AI and ethics in HRM emerges as a critical domain of scholarly and practical inquiry. Addressing these challenges requires moving beyond a purely technical focus to engage with normative considerations of fairness, accountability, transparency, and human dignity. Ethical AI in HRM is not merely a matter of compliance but a strategic imperative that shapes organisational legitimacy, employee trust, and long-term sustainability. As AI systems become increasingly embedded in workforce management, developing robust ethical frameworks and context-sensitive implementation strategies becomes essential.

This study contributes to the discourse by critically examining how AI-driven HR practices align with ethical principles while preserving their operational advantages. Drawing on perspectives from applied ethics, HRM, and information systems, it advances a more integrated understanding of the opportunities and risks associated with AI in recruitment and selection and informs the development of responsible, equitable, and human-centred practices. Although scholarly interest in AI in recruitment and selection has grown, the literature remains predominantly Western-oriented in its theoretical framing and empirical focus. Existing studies provide limited contextual sensitivity to diverse institutional and cultural settings, with minimal explicit attention to African organisational environments. This review addresses this gap by systematically synthesising evidence on ethical issues in AI-enabled HRM and highlighting the need for more contextually grounded research.

This study systematically explores ethical AI in recruitment and selection by integrating evidence across five key thematic areas: bias and discrimination in algorithmic recruitment, transparency and explainability, trust and employee dignity, organisational governance, and inclusive design strategies. It also identifies several critical gaps and areas for further research within the current literature. The review aimed to (a) investigate bias and discrimination in algorithm-based recruitment and selection; (b) analyze issues related to transparency, explainability, and accountability in AI-driven HR systems; (c) synthesize evidence on trust, perceptions of fairness, and dignity reported by various studies; (d) assess how organizational governance and ethical AI are addressed; and (e) propose strategies to reduce bias and enhance inclusive design in AI-enabled HR practices. The subsequent sections of the paper include a literature review, a discussion of methods, findings, a conclusion, and recommendations.

2. LITERATURE REVIEW

The integration of AI into human resource management has generated a rapidly expanding body of literature that draws on established organisational theories while simultaneously challenging their assumptions (Gong Fan & Bartram 2025; Singh et al. 2025). At a theoretical level, much of the scholarship is grounded in socio-technical systems theory, which posits that organisational outcomes are shaped by the joint optimisation of social and technical subsystems (Hanafizadeh and Mehra 2025). Within this framework, AI is not simply a neutral tool but an active component that restructures decision-making processes, redistributes power, and redefines accountability within HR functions. Scholars argue that algorithmic systems embed values and assumptions, making them inherently normative rather than purely technical artefacts (Selbst et al. 2019). This insight is critical in HRM, where decisions directly affect individuals' careers, livelihoods, and well-being.

Closely related is the concept of algorithmic management, which examines how organisational control is increasingly exercised through data-driven systems rather than human supervisors. Drawing on labour process theory (Zhao et al. 2026) and organisational control perspectives, this line of work highlights how AI enables new forms of surveillance, evaluation, and behavioural regulation. Empirical and conceptual studies show that algorithmic systems can intensify monitoring while reducing opportunities for employee voice and discretion (Lakshman 2026; Zhang et al. 2025), thereby reshaping the employment relationship (Kellogg et al. 2020). These developments raise important ethical questions regarding autonomy, dignity, and the balance of power between employers and employees.

Organisational justice theory further provides a robust lens for understanding employee reactions to AI-driven HR practices (Bennett and Martin 2026). The theory distinguishes between distributive justice, procedural justice, and interactional justice, all of which are implicated in algorithmic decision-making. AI systems often operate as "black boxes," limiting explainability and transparency, which in turn undermines perceptions of procedural fairness. Research indicates that when employees do not understand how decisions are made or cannot contest them, trust in organizational systems declines (Bennett and Martin 2026; Rieke and Bogen 2018; Floridi et al. 2018; Hunkenschroer and Luetge 2022; Jobin et al. 2019; Kellogg et al. 2020; Leavitt et al.

AI ETHICS IN HUMAN RESOURCE MANAGEMENT: A SYSTEMATIC REVIEW

2024; Raghavan et al. 2020; Selbst et al. 2019), even when outcomes may appear objectively consistent (Hunkenschroer and Luetge 2022). This highlights a key tension between efficiency and legitimacy in AI-enabled HRM.

At the same time, strategic HRM and resource-based perspectives frame AI as a capability that enhances organisational performance. From this viewpoint, AI-driven analytics can optimise workforce planning, improve talent acquisition, and support evidence-based decision-making, thereby contributing to sustained competitive advantage (Bennett and Martin, 2026; Hunkenschroer and Luetge, 2022; Kellogg et al., 2020; Leavitt et al., 2024). However, this instrumental perspective has been critiqued for insufficiently addressing the ethical externalities of AI adoption. It tends to prioritise efficiency gains while overlooking risks related to bias, employee well-being, and exclusion, creating an imbalance between organisational outcomes and societal responsibilities.

Empirically, research on AI in HRM reveals both its transformative potential and its ethical limitations. Studies of algorithmic hiring systems consistently demonstrate the risk of bias replication when models are trained on historically skewed data. Evidence shows that such systems can inadvertently disadvantage already marginalised groups, reinforcing patterns of inequality rather than mitigating them (Raghavan et al. 2020; Rieke and Bogen 2018). These findings directly challenge early narratives that positioned AI as a corrective to human bias, suggesting that technological systems may scale and encode discrimination in more subtle, less visible ways.

In parallel, empirical work on algorithmic management highlights the consequences of increased digital monitoring in the workplace. Employees subjected to continuous data-driven evaluation often report heightened stress, reduced autonomy, and diminished job satisfaction. These outcomes are particularly pronounced in contexts where algorithmic decision-making lacks transparency or where employees have limited avenues for recourse (Kellogg et al. 2020; Lakshman 2026). While such systems may improve measurable productivity, they also generate unintended psychosocial costs that are insufficiently captured in performance metrics.

Despite these contributions, the literature remains fragmented and exhibits several critical gaps. First, there is a lack of integration between theoretical frameworks and empirical research. While concepts such as accountability, transparency, and fairness are widely discussed, they are rarely operationalised consistently or measured. This limits the ability to systematically evaluate the ethical performance of AI systems in HRM. Second, much of the empirical evidence is context-specific, with a strong concentration in technology-intensive sectors and developed economies. This raises questions about the applicability of findings to different cultural and institutional settings, particularly in emerging economies where organisational capacities and regulatory frameworks may differ significantly.

Third, existing studies often adopt a narrow focus on specific HR functions, such as performance evaluation or recruitment, without examining how AI systems interact across the employee lifecycle. This siloed approach obscures the interconnected and cumulative nature of ethical risks. Fourth, there is a methodological imbalance in the literature, with a predominance of case studies and conceptual analyses and a relative scarcity of large-scale and longitudinal quantitative research. As a result, causal relationships between AI adoption and organisational or employee outcomes remain understudied.

These limitations collectively justify the need for a systematic review of AI and ethics in HRM. The field remains fragmented, with conceptual ambiguity, a predominance of cross-sectional and normative studies, and limited integration of empirical evidence across contexts. This makes it difficult to generate coherent insights or actionable guidance. A systematic review is therefore necessary to consolidate existing knowledge, align theoretical perspectives with available evidence, and identify clear and contextually relevant research gaps. More importantly, it provides a structured foundation for understanding how ethical principles can be operationalised within AI-driven HR practices. In doing so, the review responds to both scholarly and practical imperatives by supporting the development of more transparent, accountable, and human-centred applications of AI in HRM.

3. METHODOLOGY

3.1. Study Design

This study adopted a systematic review design guided by PRISMA protocols, with the protocol registered on the Open Science Framework to ensure transparency and reproducibility. A structured literature search was conducted across major academic databases, including Scopus, Web of Science, JSTOR, PsycINFO, and IEEE Xplore, using predefined keywords related to artificial intelligence, human resource management, and ethics. The review focused on studies published between 2019 and 2025, applying clearly defined inclusion and exclusion criteria to capture research that addresses both AI applications and ethical considerations within HR contexts. Study selection and quality appraisal were undertaken independently by multiple reviewers using established tools such as CASP and MMAT, with disagreements resolved through consensus. Data were then systematically extracted and synthesised to identify key ethical themes, methodological trends, and research gaps, with particular attention to contextual and underrepresented perspectives.

3.2. Data Collection

The PRISMA (Preferred Reporting Items for Systematic Reviews and Meta-Analyses) approach was used for data collection. This is because researchers have suggested the approach as effective (Page et al. 2021; Mishra and Mishra 2023). The most commonly agreed-upon format for presenting evidence in systematic reviews and meta-analyses is PRISMA (Dwinggo et al. 2023). Peer-reviewed journal articles and conference proceedings were taken into consideration in this review. The choices of the articles were made according to the following four criteria: academic journals or conferences, publication dates between January 2019 and January 2025, English language, and the substantive focus on both the application of AI and the ethical implications of the use of AI in HR. The review focused on studies from 2019 to 2025, a period marked by rapid growth in AI-driven HR tools following advances in generative AI and machine learning. Earlier research tended to concentrate on conceptual aspects of AI automation or HR analytics, without addressing the ethical issues of recent AI-based HR systems. Limiting the review to this timeframe ensured that the synthesised evidence highlighted governance challenges, current technologies, and organisational practices. It excluded non-peer-reviewed conference proceedings, reports, books, editorials, opinion pieces, commentaries, grey literature, non-English publications, studies describing AI ethics without HR applications, and studies examining HR without AI components. The databases were searched in Web of Science (Clarivate Analytics), JSTOR, PsycINFO (EBSCOhost), IEEE Xplore, and ProQuest. To maximise retrieval while maintaining relevance, a comprehensive Boolean search strategy was developed using combinations of keywords and synonyms connected by the operators OR and AND. The search strategy combined three principal concepts: (1) artificial intelligence technologies, (2) human resource management functions, and (3) ethical considerations. A representative search string was: "machine learning" OR "artificial intelligence" OR AI OR "predictive analytics" OR "algorithmic decision-making" OR/ AND ("human resource management" OR HRM OR recruitment OR hiring OR selection OR "talent management" OR "performance appraisal" OR "employee management") AND (fairness OR ethics OR transparency OR bias OR discrimination OR explainability OR accountability OR employee trust OR governance). The search syntax was adapted where necessary to accommodate the search functionalities and indexing requirements of each database. Additional manual screening of reference lists from eligible studies was undertaken to identify any relevant articles that may not have been captured through the electronic database searches. EndNote was used for reference management, while deduplication and collaborative screening were conducted using Rayyan. A critical appraisal was conducted using the CASP checklists and the Mixed Methods Appraisal Tool (MMAT).

The study selection process followed PRISMA 2020 guidelines, with studies screened through a structured multi-stage procedure to ensure relevance and quality. Following database retrieval, records were progressively excluded based on predefined criteria, including lack of direct relevance to both AI applications and ethical considerations in HRM, as well as failure to meet quality appraisal standards. Although the initial database search retrieved 347 records, applying specific eligibility criteria significantly narrowed down the eligible studies. This review focused on peer-reviewed, empirical research that directly examined both AI applications and ethical concerns within HRM. Many of the retrieved publications discussed AI without an ethical focus, addressed AI ethics outside HRM, or were conceptual without empirical data. Figure 1 illustrates the article screening process used in this study.

AI ETHICS IN HUMAN RESOURCE MANAGEMENT: A SYSTEMATIC REVIEW

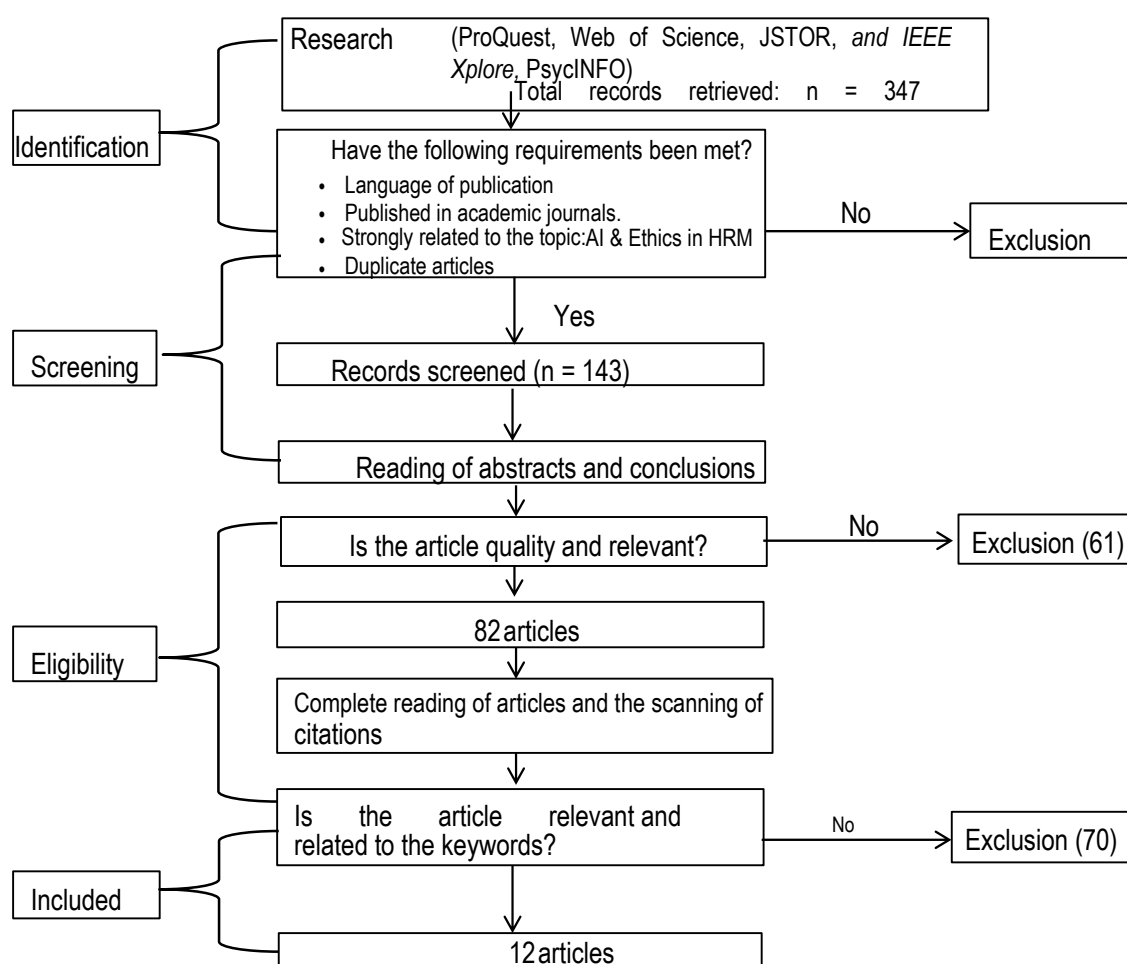


FIGURE 1 - PRISMA 2020 FLOW DIAGRAM

As a result, the final selection reflects the current developmental stage of this interdisciplinary field rather than search limitations. Similar systematic reviews in emerging research areas often find small evidence bases when strict inclusion criteria are used. Including 12 studies indicates the rigorous application of the review protocol rather than a lack of relevant research. AI ethics in HRM is a specialised area at the crossroads of artificial intelligence, ethics, and human resource management. Many initial records covered only one of these aspects and did not meet the criteria. Thematic analysis showed consistent focus on core ethical issues: algorithmic bias, transparency, governance, accountability, bias reduction, and trust, with few new ideas appearing later. This pattern suggests thematic saturation for the review's goals. While more studies could add context, they are unlikely to change the main ethical themes. Still, the small evidence base highlights the need for further empirical research across various organisations and regions. Research explicitly addressing both AI implementation and ethical implications within HR remains limited. Despite the relatively small number of included studies, the sample provides sufficient depth to support a focused thematic synthesis and to identify consistent patterns and gaps within the literature.

3.3. Data Analysis

Content analysis was employed to systematically examine the extracted data using both descriptive and thematic approaches. First, a descriptive analysis was conducted to summarise the characteristics of the included studies, including study design, context, and key findings, enabling comparison across the evidence base. This was followed by a thematic analysis, in which an interpretive approach was used to identify and synthesise recurring patterns within the data. Themes were developed through a hybrid coding process that combined deductive coding, guided by predefined analytical categories derived from the review protocol, and inductive coding, allowing new themes to emerge from the data. This iterative process resulted in the identification of five overarching thematic areas.

AI ETHICS IN HUMAN RESOURCE MANAGEMENT: A SYSTEMATIC REVIEW

3.4. Limitations of the Study

This review has several limitations. First, it was restricted to peer-reviewed, English-language studies published between 2019 and 2025. While this ensured the inclusion of contemporary evidence reflecting the rapid evolution of AI ethics in HRM, relevant studies published in other languages or prior to 2019 may not have been captured. Second, the review focused on five major multidisciplinary databases that provide extensive coverage of management, information systems and social science research. Although additional sources may contain relevant studies, the selected databases were considered sufficient to identify the principal body of peer-reviewed evidence in this emerging field. Finally, only 12 studies met the predefined inclusion criteria. Rather than reflecting limitations in the search strategy, this small evidence base highlights the emerging nature of empirical research on AI ethics in HRM. Despite the limited number of eligible studies, consistent patterns were observed across the evidence, with recurring themes relating to bias, transparency, trust, governance and ethical AI implementation. These findings suggest that the review provides a robust synthesis of the current state of knowledge while identifying important directions for future research. Accordingly, the principal contribution of this review is not the size of the evidence base but the integration and critical synthesis of the available empirical literature into a coherent understanding of the ethical challenges, their interrelationships and the governance mechanisms required for responsible AI implementation in HRM.

4. FINDINGS

The findings reveal five main themes in ethical AI within HRM (Table 1)

TABLE 1 - SUMMARY OF KEY FINDINGS

Theme	Description	Reviewed Studies	Key Gap Identified
Bias and discrimination in algorithm-based recruitment and selection	AI systems reproduce and scale historical inequalities embedded in training data, leading to discriminatory hiring and evaluation outcomes	Tambe et al. 2019; Newman et al. 2020; Muralidharan et.al. 2024; Arora et.al. 2021; Budhwar et al. 2022; Hunkenschroer and Luetge 2022; Madanchian, 2024, Raghavan et.al. 2020; Rao and Zhao 2025; Varma et al. 2023	Limited cross-cultural validation; insufficient focus on non-Western and African labor markets. Heavy reliance on conceptual and cross-sectional studies; limited longitudinal and causal evidence; insufficient mixed-methods validation. Africa remains understudied.
Transparency, explainability, and accountability	Algorithmic opacity limits understanding of decision processes and complicates accountability, raising ethical and legal concerns	Muralidharan et.al. 2024; Hunkenschroer and Luetge 2022; Budhwar et al. 2022; Arora et.al. 2021	Lack of empirical validation of explainable AI frameworks; weak governance structures in practice. Limited empirical testing of explainable AI frameworks; lack of standardized measures for transparency and accountability. Weak evidence from non-Western regulatory environments; limited insight into how transparency is operationalized in African organizations
Trust, fairness perceptions, and employee dignity	Employee trust depends on perceived fairness, transparency, and procedural justice, while AI may undermine autonomy and meaningful work	Singh et.al. 2025; Leavitt et. al. 2024; Bankins and Formosa 2023; Kellogg et al. 2020	Limited cultural diversity in samples; absence of African perspectives on trust, dignity, and workplace norms. Limited longitudinal evidence; insufficient exploration of employee perceptions in diverse contexts. Predominantly cross-sectional and perception-based studies;
Organizational governance and ethical AI systems	Effective governance mechanisms, oversight structures, and AI literacy are essential but remain underdeveloped in organizations	Tambe et al. 2019; Budhwar et al. 2022; Hunkenschroer and Luetge 2022; Madanchian,2024, Varma et al. 2023; Muralidharan et. al., 2024; Rao and Zhao 2025; Bankins and Formosa 2023	Governance gaps, especially in developing and African contexts; lack of standardized and enforceable frameworks. Dominance of conceptual and normative analyses; limited empirical validation of governance models; lack of comparative studies
Strategies for reducing bias and promoting inclusive design	Technical and organizational interventions such as debiasing techniques, participatory design, and auditing can mitigate ethical risks	Muralidharan et.al. 2024; Hunkenschroer and Luetge 2022; Rao and Zhao 2025	Limited integration between technical and organizational solutions; lack of real-world implementation and scalability evidence. Minimal testing of interventions in diverse socio-cultural settings; lack of adaptation to African organizational and institutional contexts

AI ETHICS IN HUMAN RESOURCE MANAGEMENT: A SYSTEMATIC REVIEW

These include biases and discrimination in algorithmic hiring, where AI tends to mirror existing inequalities; issues with transparency, explainability, and accountability due to opaque decision-making; concerns about trust, fairness, and employee dignity that influence AI acceptance; weak organisational governance and underdeveloped ethical structures; and new approaches to reduce bias and support inclusive design. Across all themes, there is a persistent lack of long-term, diverse contextual evidence, especially from non-Western and African regions.

4.1. Discussions

The integration of artificial intelligence into human resource management is accelerating, yet the literature reveals persistent and systemic ethical concerns across key HR functions. Evidence shows that algorithmic bias remains widespread, with limited cross-cultural validation, a notable absence of research in African contexts, and a heavy reliance on conceptual and cross-sectional studies. Similarly, issues of transparency and accountability are constrained by weak empirical validation, a lack of standardised measures, and a limited understanding of how these principles operate outside Western settings. Research on trust, fairness, and employee dignity is also methodologically narrow and contextually limited, with insufficient longitudinal evidence and minimal representation of diverse cultural perspectives. Furthermore, governance frameworks and bias mitigation strategies remain underdeveloped, largely theoretical, and rarely tested in real-world or non-Western environments, underscoring a critical need for more empirically grounded, context-sensitive research.

4.1.1 Bias and discrimination in algorithm-based recruitment and selection

Algorithmic bias was the most common and frequently discussed ethical concern across the reviewed literature, and it was notably highlighted in 10 of the 12 studies considered. It is widely accepted that AI systems, trained on historical HR data, tend to reproduce, and often reinforce, existing patterns of discrimination, particularly concerning gender, race, ethnicity, and socio-economic background. Tambe et al. (2019) highlight significant data limitations in HR contexts, including small sample sizes, fragmented datasets, and non-standardised data, which can increase the risk of biased or unreliable AI-driven decisions if not carefully managed. In addition, Hunkenschroer and Luetge (2022) provide a comprehensive review of ethical challenges in AI-enabled recruitment, emphasising bias, transparency, and the potential loss of human-centred decision-making, and call for stronger ethical governance and future research. Rao and Zhao (2025) use a case-based analysis of AI-enabled hiring to demonstrate how algorithmic decisions can be influenced by features that are indirectly correlated with protected attributes. Even when sensitive variables are excluded, the model may rely on proxy indicators embedded in résumé data, resulting in uneven outcomes across applicant groups. This illustrates the risk of proxy discrimination, in which bias arises from the structure of the data and feature selection rather than from explicit intent. Their findings highlight that eliminating protected variables alone is insufficient to ensure fairness. Instead, they emphasise the need for systematic evaluation of fairness and careful model design to mitigate such risks.

Newman et al. (2020) show that reducing statistical bias in algorithmic HR decisions does not necessarily translate into perceived fairness, as employees' reactions are strongly shaped by procedural justice. Their findings highlight the importance of factors such as voice, transparency, and individualised consideration in shaping perceptions of fairness, with implications for the design of AI-enabled HR systems. Varma et al. (2023) adopt a critical ethical lens to examine AI in HRM, highlighting how its deployment can reinforce existing power asymmetries and inequalities between employers and employees. They caution that, without appropriate governance, AI-enabled practices may disproportionately affect already disadvantaged groups.

Algorithmic bias persists not simply because AI systems inherit biased historical data, but because bias is embedded throughout the broader socio-technical HR ecosystem. Several studies suggest that organisations frequently treat fairness as a technical problem to be solved through data cleaning or model refinement, while giving less attention to organisational practices, governance structures, and institutional inequalities that shape how AI systems are developed and deployed (Hunkenschroer & Luetge, 2022; Raghavan et al., 2020; Selbst et al., 2019). Consequently, even highly advanced models might still generate biased results if recruitment criteria, performance metrics, or past employment data already embody systemic discrimination. This is why

AI ETHICS IN HUMAN RESOURCE MANAGEMENT: A SYSTEMATIC REVIEW

simply eliminating protected attributes like gender or race often fails to prevent unfair outcomes, as proxy variables and organisational decision rules continue to carry forward historical inequalities into the predictions.

The reviewed studies also show that algorithmic bias does not occur in isolation but interacts with other ethical challenges identified in this review. Limited transparency makes it difficult for applicants and employees to detect or contest discriminatory decisions, while weak governance reduces organisational accountability for biased outcomes. These conditions, in turn, undermine employee trust and perceptions of procedural justice, creating a reinforcing cycle in which opaque systems conceal bias, weak oversight allows it to persist, and declining trust erodes organisational legitimacy. Rather than representing separate ethical concerns, bias, transparency, governance, and trust emerge as interconnected dimensions of responsible AI implementation in HRM.

The evidence further indicates that no single intervention is sufficient to mitigate algorithmic bias. Technical approaches, including fairness-aware machine learning, bias detection, data rebalancing, and algorithmic auditing, improve model performance but cannot address biases originating from organisational policies or historical labour market inequalities (Muralidharan et al., 2024; Rao & Zhao, 2025). Conversely, governance interventions such as human oversight, ethical review processes, explainability requirements, and interdisciplinary collaboration strengthen accountability but are less effective without robust technical safeguards (Budhwar et al., 2022; Tambe et al., 2019). Collectively, the literature suggests that the most effective approach combines technical fairness mechanisms with organisational governance, continuous monitoring, and meaningful human oversight throughout the AI lifecycle. However, empirical evidence evaluating these integrated interventions remains limited, particularly in non-Western and African contexts, highlighting an important direction for future research.

4.1.2 Transparency, Explainability, and Accountability

A key finding of this review is that transparency and explainability remain major ethical challenges in AI-driven HR systems. Many studies describe these systems as opaque, limiting understanding of how decisions in recruitment, evaluation, and promotion are made. This lack of clarity undermines procedural fairness and makes it difficult for organisations to justify outcomes. The literature consistently positions transparency as essential for ethical AI. Muralidharan et al. (2024) emphasise the importance of explainable AI for enhancing transparency and interpretability in HR decision-making, particularly in algorithm-driven recruitment and evaluation systems. Similarly, Hunkenschroer and Luetge (2022) identify transparency as a central ethical concern, closely associated with issues of accountability and fairness in AI-enabled HR practices. Evidence from Budhwar et al. (2022) and Arora et al. (2021) further indicates that organisations continue to face notable transparency challenges, with governance structures and oversight mechanisms often lagging the rapid pace of AI adoption. Collectively, these studies highlight a persistent gap between the technical deployment of AI systems and the organisational capacity to ensure their ethical and transparent use.

A critical issue is the resulting accountability gap, where responsibility for decisions becomes unclear when both algorithms and humans are involved. This weakens oversight and increases ethical and legal risks. Additionally, lack of transparency negatively affects employee trust, as individuals are less likely to accept decisions they cannot understand or challenge. Overall, the findings indicate that while transparency is widely recognised as essential, its practical implementation remains limited, particularly in non-Western contexts. This highlights the need for stronger, context-sensitive governance frameworks that ensure AI systems are both explainable and accountable.

Beyond recognising transparency as a valuable principle, the literature indicates that implementing it is difficult because transparency involves more than just technical aspects; it also encompasses organisational and governance challenges. Explainable AI methods can enhance comprehension of algorithmic results, but many organisations lack the necessary policies, expertise, and accountability frameworks to turn technical explanations into meaningful decisions for HR professionals and employees (Hunkenschroer & Luetge, 2022; Muralidharan et al., 2024). Additionally, many commercial AI systems function as proprietary "black boxes," restricting organisations from thoroughly examining decision processes or clearly communicating them to employees. As a result, transparency often remains an ideal rather than a practical implementation.

AI ETHICS IN HUMAN RESOURCE MANAGEMENT: A SYSTEMATIC REVIEW

The results also show that transparency is deeply linked to the other ethical aspects discussed in this review. When explainability is limited, it becomes harder to identify algorithmic bias, weakens organisational accountability, and erodes employees' trust in the fairness and contestability of HR decisions. Conversely, increased transparency enhances governance by clarifying roles, supporting independent audits, and allowing employees to question or appeal algorithmic choices. These connections imply that transparency is a key enabler for fairness, accountability, and trust, rather than just a standalone ethical standard. Without sufficient transparency, organisations cannot properly assess if AI systems are fair or ensure compliance with ethical and regulatory standards.

The literature indicates that the most effective interventions often combine technical and organisational strategies. Technical approaches, such as explainable AI models, documentation of development processes, and algorithmic audits, enhance AI interpretability but are less effective without supportive governance measures. Organisational strategies like human oversight, clear accountability structures, employee communication, and regular ethical reviews tend to be more successful when paired with technical explainability tools (Budhwar et al., 2022; Tambe et al., 2019; Muralidharan et al., 2024). Nonetheless, there are few empirical studies systematically assessing these combined methods, especially in developing economies and African organisational settings. This gap hampers understanding of how transparency can be reliably implemented across diverse institutional and cultural contexts.

4.1.3 Trust, fairness perceptions, and employee dignity

The empirical studies conducted found that transparency and perceived fairness were the two most important factors influencing trust in AI-HR systems. Newman et al. (2020) demonstrate that outcome fairness and procedural fairness are distinct dimensions that independently shape employee perceptions of fairness in HR decisions. Their findings show that even statistically optimised, algorithm-driven decisions may be perceived as unfair when they do not allow for employee voice or fail to account for individual circumstances. The study highlights that perceptions of fairness are not determined solely by the accuracy or objectivity of outcomes but are strongly influenced by how decisions are made and communicated. In this regard, procedural justice mechanisms such as opportunities for employee input, clear communication of decision criteria, and accessible avenues for appeal play a critical role in shaping acceptance of AI-enabled HR systems.

Bankins and Formosa (2023) examine the ethical implications of AI for meaningful work, highlighting how its use may affect employee autonomy, skill development, and sense of purpose. They argue that the increasing reliance on algorithmic decision-making can reshape the nature of work, with potential implications for how employees experience meaning in their roles. In particular, the delegation of judgment to AI systems may reduce opportunities for discretion and learning, thereby limiting employees' ability to derive fulfilment from their work. This raises broader concerns about the erosion of intrinsic motivation and the quality of work experience in AI-mediated environments. Madanchian (2024) complements this perspective by emphasising the need to balance AI's efficiency gains with broader organisational and human considerations.

The studies reviewed indicate that employee trust in AI-powered HR systems is more influenced by perceptions of transparency, fairness, and respect for dignity than by the accuracy of algorithms. Trust develops as a relational outcome through organisational processes, not solely from technological performance. Even if AI systems enhance consistency and efficiency, employees might stay sceptical if they cannot understand decision-making processes, feel their personal situations are ignored, or believe they lack channels to question or appeal decisions (Newman et al., 2020; Leavitt et al., 2024). These findings suggest that building trust is challenging because ethical acceptance relies equally on procedural justice, organisational communication, and algorithmic accuracy.

The evidence further shows that trust is intimately linked to the wider ethical issues discussed in this review. Perceptions of algorithmic bias erode confidence in recruitment and promotion processes, while a lack of transparency hampers employees' ability to understand or challenge organisational decisions. Likewise, poor governance structures diminish confidence that organisations can effectively oversee AI systems or address unfair outcomes. Therefore, trust, fairness, transparency, and governance should be viewed as interconnected elements that reinforce each other rather than separate ethical problems. If an organisation neglects one aspect, it is likely to weaken the others, which can reduce employee support for AI-based HR practices and

AI ETHICS IN HUMAN RESOURCE MANAGEMENT: A SYSTEMATIC REVIEW

potentially damage organisational legitimacy. The literature suggests that the most effective ways to strengthen trust involve more than just technical improvements; they also include organisational practices that promote procedural justice. Key elements such as human oversight, employee participation and appeal opportunities, transparent communication about AI-assisted decisions, and clear accountability mechanisms are consistently more effective than focusing solely on algorithmic accuracy or predictive performance. (Leavitt et al., 2024; Newman et al., 2020; Bankins & Formosa, 2023). Nonetheless, there is limited empirical data on the long-term success of these interventions, especially when considering varied cultural and institutional settings. Future studies should explore how trust evolves over time and whether governance approaches tailored to specific contexts can maintain employee confidence in AI-driven HR systems across different organisations.

4.1.4 Organisational governance and ethical AI systems

Nine out of twelve studies identified governance deficits or addressed aspects of ethical governance frameworks. Arora et al. (2021) highlight the emerging ethical and governance challenges associated with the adoption of AI in HR, noting the absence of clearly defined frameworks tailored to its human and organisational implications. Similarly, Budhwar et al. (2022) provide empirical evidence that organisations, particularly in high-tech contexts, face uncertainty and inconsistencies in governing AI-enabled HR practices, with ethical considerations often lagging technological adoption. Their findings suggest that while AI tools are increasingly deployed, corresponding governance structures remain underdeveloped and fragmented. This misalignment raises concerns about accountability, oversight, and organisations' ability to effectively manage the ethical risks associated with AI-driven decision-making.

Hunkenschroer and Luetge (2022) provide one of the more comprehensive ethical analyses in the literature, identifying key dimensions such as fairness, autonomy, transparency, accountability, and data privacy in AI-enabled HR practices. Their work draws on multiple ethical perspectives to frame the complex normative challenges associated with algorithmic decision-making, although it does not formalise these into a single unified framework. Rather, it offers a structured foundation for understanding the ethical tensions that arise when efficiency-driven technologies intersect with human-centred HR practices. Complementing this, Tambe et al. (2019) emphasise the importance of interdisciplinary collaboration among HR professionals, data scientists, and other stakeholders to address the ethical and governance challenges of AI in HR, highlighting the need for integrated expertise to manage its organisational implications. They further argue that effective implementation of AI in HR requires alignment between technical capabilities and organisational processes, particularly regarding data quality, decision-making structures, and accountability mechanisms. Together, these contributions underscore the importance of combining ethical reflection with practical governance approaches in advancing responsible AI adoption in HRM.

The review highlights a notable gap in the contextual application of AI ethics frameworks within African and other non-Western organisational environments. Much of the existing literature is grounded in Western regulatory contexts and philosophical traditions, which may not fully capture the diversity of institutional arrangements, cultural norms, and labour market dynamics across different regions. Consequently, there remains a limited understanding of how these frameworks translate to contexts with differing governance capacities and socio-cultural foundations. Emerging perspectives from African philosophical traditions, such as Ubuntu, with its emphasis on relational identity, communal dignity, and collective responsibility, offer potentially valuable insights for developing more contextually grounded and inclusive approaches to AI ethics in HRM, though these remain underexplored in current scholarship.

The evidence indicates that governance gaps remain because the adoption of AI in organisations has often advanced faster than the creation of proper ethical oversight mechanisms. Although many organisations have invested in AI to boost efficiency and decision-making, fewer have set up formal governance frameworks that explicitly assign accountability, outline ethical responsibilities, and oversee the entire AI lifecycle (Budhwar et al., 2022; Hunkenschroer & Luetge, 2022). This imbalance is worsened by low AI literacy among HR professionals, fragmented regulations, and the lack of standard governance frameworks adaptable to various organizational and cultural settings. Consequently, ethical governance often remains reactive instead of being integrated into the entire AI lifecycle, from design to monitoring.

AI ETHICS IN HUMAN RESOURCE MANAGEMENT: A SYSTEMATIC REVIEW

The findings also show that governance acts as the central mechanism that connects the other ethical issues discussed in this review. Strong governance offers the necessary structures to identify and reduce algorithmic bias, promote transparency through documentation and explanation, and build employee trust by defining accountability and offering oversight and appeal options. On the other hand, weak governance allows biased systems to go unchecked, hampers transparency, and diminishes confidence in the ability to challenge or correct AI-driven decisions. Therefore, governance is not a separate ethical issue but a foundation for ensuring fairness, transparency, accountability, and respect for employee dignity in AI-driven HR practices.

Across the reviewed studies, the most effective governance approaches encompass a combination of organisational, technical, and regulatory measures. Organisations that incorporate multidisciplinary oversight involving data scientists, human resources professionals, senior leadership, and legal experts are better equipped to identify ethical risks throughout the AI lifecycle than those relying solely on technical controls (Tambe et al., 2019; Budhwar et al., 2022). Likewise, governance measures like independent algorithm audits, ongoing monitoring, well-defined accountability frameworks, AI ethics policies, and human oversight are seen as complementary rather than competing approaches. Nonetheless, there is limited empirical evidence on the effectiveness of these integrated governance strategies, especially within developing economies and African organisational contexts. This underscores the importance of conducting comparative and longitudinal studies to understand how governance frameworks impact the ethical performance and long-term legitimacy of AI-driven HR systems.

4.1.5 Strategies for reducing bias and promoting inclusive design.

Several studies went a step further, focusing not on diagnosis but on specific measures for the ethical use of AI in HR. Muralidharan et al. (2024) emphasise the need to balance predictive efficiency with fairness in applying machine learning to HR practices, advocating the integration of bias detection and fairness evaluation throughout the model development lifecycle. They also highlight the importance of human oversight in supporting transparent and accountable decision-making, noting that ethical considerations must be embedded within both technical design and organisational processes. Their approach underscores that fairness cannot be treated as an afterthought but must be systematically incorporated from data preparation through to model evaluation. Similarly, Rao and Zhao (2025) examine the use of machine learning as an analytical tool for evaluating hiring decisions, demonstrating how bias can be identified and assessed within recruitment processes. They recommend incorporating fairness evaluation, improving data preparation, and systematically assessing model outputs, while using multiple models to validate decision factors and support more informed and equitable HR decision-making. Together, these studies highlight the importance of combining technical rigour with ethical oversight in the responsible deployment of AI in HRM.

Budhwar et al. (2022) and Arora et al. (2021) both defined AI literacy among HR professionals as a significant yet underdeveloped organisational aptitude. Ethical governance demands HR professionals who can critically scrutinise AI tools, interpret results, and identify potential biases. They must also be able to share decisions with the impacted employees with the necessary level of transparency and empathy. HR staff training on AI was described as an ethical necessity rather than a technical one.

The evidence suggests that reducing algorithmic bias requires shifting from isolated technical solutions to comprehensive socio-technical strategies that tackle both algorithmic performance and organisational decision-making. Bias often remains because interventions are typically focused only on a single stage of the AI lifecycle, such as during model development, while aspects like data governance, organisational policies, implementation techniques, and continuous evaluation receive less attention. As a result, improving model fairness alone does not guarantee fair HR outcomes if organisational processes still perpetuate historical inequalities or if AI-generated suggestions are accepted without critical human oversight.

The findings additionally indicate that technical and organisational interventions function synergistically rather than in competition. Technical measures, such as fairness-aware machine learning, algorithmic auditing, bias detection, and enhanced data preparation, serve to improve the capacity to identify and mitigate discriminatory outcomes (Muralidharan et al., 2024; Rao & Zhao, 2025). However, these strategies work best when paired with organisational efforts like AI literacy for HR professionals, participatory system design, multidisciplinary oversight, transparent communication, and active human participation in critical employment decisions.

AI ETHICS IN HUMAN RESOURCE MANAGEMENT: A SYSTEMATIC REVIEW

(Budhwar et al., 2022; Arora et al., 2021). This interaction emphasises that responsible AI deployment relies not just on enhancing algorithms but also on building institutional capacity to oversee their ethical use.

Across the reviewed studies, no single intervention reliably eliminates bias. Instead, the best-supported approach involves integrated governance frameworks that include technical safeguards, ongoing monitoring, ethical oversight, and employee participation throughout the AI lifecycle. However, empirical evidence for these combined strategies remains scarce, mostly based on conceptual analysis or short-term case studies. Additionally, there is limited data on how bias mitigation efforts perform across different cultural, institutional, and regulatory contexts, especially within African organisations. Future research should focus on assessing the long-term effectiveness of these combined technical and organisational strategies in real-world HR environments to develop evidence-based best practices for ethical AI deployment.

5. CONCLUSION AND RECOMMENDATIONS

5.1 Conclusion

The systematic review posits that while the integration of artificial intelligence into human resource management offers measurable efficiency gains and the potential to reduce certain forms of human bias, it also introduces recurring ethical challenges. Across the 12 studies reviewed, there is consistent evidence of concerns about bias, limited transparency, evolving accountability structures, and implications for employee trust and dignity. These patterns suggest that such challenges are not isolated but reflect broader issues associated with the current implementation of AI in HRM. A central tension emerging from the literature is the balance between efficiency and ethical considerations. Although AI-driven HR systems enhance speed, scalability, and data-driven decision-making, these benefits may be accompanied by limitations in transparency, fairness, and human-centred practices. The evidence indicates that addressing these challenges requires more than technical adjustments; it calls for strengthened organisational governance, clearer accountability mechanisms, and the integration of ethical considerations into system design and deployment.

The review also identifies a notable gap in contextually grounded research, particularly within African organisational settings. The predominance of studies situated in Western contexts limits the generalizability of existing frameworks across diverse socio-cultural environments. This highlights the need for greater attention to context-sensitive approaches that reflect varying institutional and cultural realities. The findings underscore the importance of developing more integrated and context-aware approaches to AI ethics in HRM. For practitioners, policymakers, and researchers, the priority is to ensure that AI-enabled HR systems support not only organisational performance but also fairness, transparency, and employee well-being across different contexts.

5.2 Recommendations

The analysis indicates that AI-powered HR systems pose ethical risks that require deliberate, ongoing organisational oversight, particularly considering the governance, transparency, and bias-related gaps identified in the literature. Organisations should therefore develop context-specific AI ethics policies that clearly define accountability structures, incorporate regular bias audits, and establish formal procedures for employee review and appeal. To address the limited operationalisation of ethical principles and the identified gaps in organisational capacity, AI literacy should be promoted as a core competency among HR professionals. In addition, participatory design approaches that involve employees and relevant stakeholders are essential for addressing concerns about trust, procedural fairness, and contextual sensitivity. The lack of robust oversight and standardised evaluation mechanisms further suggests the need for regular independent algorithm audits, alongside the integration of human-in-the-loop systems to maintain accountability and oversight in high-stakes HR decisions.

The review also identifies a significant gap in contextually grounded research, particularly within African organisational environments. Accordingly, future studies should prioritise both qualitative and quantitative investigations into how AI ethics are experienced and implemented across diverse institutional settings. Such research should engage with indigenous ethical perspectives, including Ubuntu, to support the development of more context-sensitive and culturally relevant governance

AI ETHICS IN HUMAN RESOURCE MANAGEMENT: A SYSTEMATIC REVIEW

frameworks. Policymakers are encouraged to design AI governance approaches that reflect local institutional capacities and socio-cultural dynamics rather than relying solely on externally developed models.

Finally, the literature reveals a narrow focus on recruitment-related applications and a broader methodological imbalance characterised by cross-sectional and conceptual approaches. Future research should therefore extend to other HR domains, including performance management, compensation, employee monitoring, and career development, while also prioritising longitudinal designs that examine the sustained effects of AI adoption on employee outcomes such as well-being, trust, and organisational commitment. Such efforts would contribute to a more robust and empirically grounded understanding of ethical AI in HRM.

REFERENCES

- Arora, M., Prakash, A., Mittal, A., & Singh, S. (2021). HR analytics and artificial intelligence: Transforming human resource management. In *Proceedings of the 2021 International Conference on Decision Aid Sciences and Application (DASA)* (pp. 288–293). IEEE.
- Bankins, S., & Formosa, P. (2023). The ethical implications of artificial intelligence (AI) for meaningful work. *Journal of Business Ethics*, 185(4), 725–740.
- Bastida, M., Vaquero García, A., Vázquez Taín, M. Á., & del Río Araujo, M. (2025). From automation to augmentation: Human resources' journey with artificial intelligence. *Journal of Industrial Information Integration*, 46, 100872. <https://doi.org/10.1016/j.jii.2025.100872>.
- Bennett, N., and Martin, C. L. (2026). AI as a talent management tool: An organisational justice perspective. *Business Horizons*, 69(2), 241–252. <https://doi.org/10.1016/j.bushor.2025.03.005>.
- Budhwar, P., Malik, A., De Silva, M., & Thedosiou, M. (2022). Artificial intelligence human resource management: Ethical concerns among middle-level managers in high-tech corporations. *International journal of human resource management*. Vol. 33 (6), 1065-1097.
- Cambridge Consultants. (2025). *AI and moral ethics*. Retrieved on December 11, 2025, from <https://www.thecambridgeconsultancy.com/wp-content/uploads/2025/09/Ai-and-Moral-Ethics.pdf>.
- Dastin, J. (2018, October 11). *Amazon scraps secret AI recruiting tool that showed bias against women*. Reuters. Retrieved on December 9, 2025, from <https://www.reuters.com/article/world/insight-amazon-scraps-secret-ai-recruiting-tool-that-showed-bias-against-women-idUSKCN1MK0AG>.
- Dwinggo Samala, A., Usmeldi, T. A., Bojić, L., Indarta, Y., Tsoy, D., Denden, M., ... Parma Dewi, I. (2023). Metaverse technologies in education: A systematic literature review using PRISMA. *International Journal of Emerging Technologies in Learning (iJET)*, 18(5), 1–12.
- European Commission. (2019). Ethics guidelines for trustworthy AI. Retrieved on October 30, 2025, from <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>.
- Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., and Vayena, E. (2018). AI4People: An ethical framework for a good AI society. *Minds and Machines*, 28(4), 689–707. <https://doi.org/10.1007/s11023-018-9482-5>.
- Gong, Q., Fan, D., & Bartram, T. (2025). Integrating artificial intelligence and human resource management: a review and future research agenda. *The International Journal of Human Resource Management*, 36(1), 103–141. <https://doi.org/10.1080/09585192.2024.2440065>.
- Hanafizadeh P, Mehrasa S (2025), "Governance system design model in platform ecosystems by a socio-technical systems theory". *Digital Policy, Regulation and Governance*, 28 (4): 432–457. <https://doi.org/10.1108/DPRG-04-2025-0105>.

AI ETHICS IN HUMAN RESOURCE MANAGEMENT: A SYSTEMATIC REVIEW

- Hunkenschroer, A. L., and Luetge, C. (2022). Ethics of AI-enabled recruiting and selection: A review and research agenda. *Journal of Business Ethics*, 178, 977–993. <https://doi.org/10.1007/s10551-022-05049-6>.
- Jobin, A., Ienca, M., and Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1, 389–399. <https://doi.org/10.1038/s42256-019-0088-2>.
- Kellogg, K. C., Valentine, M. A., and 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>.
- Leavitt, K., Barnes, C. M., and Shapiro, D. L. (2024). The role of human managers within algorithmic performance management systems: A process model of employee trust in managers through reflexivity. *Academy of Management Review*, 50(4). <https://doi.org/10.5465/amr.2022.0058>.
- Madanchian, M. (2024). From Recruitment to Retention: AI Tools for Human Resource Decision-Making. *Applied Sciences*, 14(24), 11750. <https://doi.org/10.3390/app142411750>.
- Mishra, V., & Mishra, M. P. (2023). PRISMA for review of management literature: Method, merits, and limitations. In *Advancing methodologies of conducting literature review in the management domain* (pp. 125–136).
- Muralidharan, V., Ravi, B., Pachar, S., Jain, N., Chakraborty, N., & Lenin, D. S. (2024). Ethical considerations in applying machine learning to HR practices: Balancing efficiency with fairness. In *the 7th International Conference on Contemporary Computing and Informatics (IC3I)* (pp. 1392–1398). IEEE. <https://doi.org/10.1109/IC3I61595.2024.10828676>.
- 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. *Organisational Behaviour and Human Decision Processes*, 160, 149–167.
- Organisation for Economic Co-operation and Development (OECD). (2019). OECD principles on artificial intelligence. Retrieved on March 10, 2026, from <https://oecd.ai/en/ai-principles>.
- Page, M. J., McKenzie, J. E., Bossuyt, P. M., Boutron, I., Hoffmann, T. C., Mulrow, C. D., ... Moher, D. (2021). Updating guidance for reporting systematic reviews: Development of the PRISMA 2020 statement. *Journal of Clinical Epidemiology*, 134, 103–112.
- Panarese, P., Grasso, M. M., & Solinas, C. (2025). Algorithmic bias, fairness, and inclusivity: A multilevel framework for justice-oriented AI. *AI & Society*. <https://doi.org/10.1007/s00146-025-02451-2>.
- Pliszka, B. (2025). Artificial intelligence in human resources management: A momentary trend or necessity? Scientific Papers of Silesian University of Technology. *Organisation and Management Series*, 229, 449–462. <https://doi.org/10.29119/1641-3466.2025.229.25>.
- Rabenu, E., & Baruch, Y. (2025). Cyborging HRM theory: From evolution to revolution – The challenges and trajectories of AI for the future role of HRM. *Personnel Review*, 54(1), 174–198. <https://doi.org/10.1108/PR-02-2024-0111>.
- Raghavan, M., Barocas, S., Kleinberg, J., and Levy, K. (2020). Mitigating bias in algorithmic hiring: Evaluating claims and practices. In *Proceedings of the ACM Conference on Fairness, Accountability, and Transparency*. <https://doi.org/10.1145/3351095.3372828>.
- Rao, S., & Zhao, T. (2025). Ethical AI in HR: A Case Study of Tech Hiring. *Journal of Computer Information Systems*. <https://doi.org/10.1080/08874417.2024.2446954>.
- Rieke, A. and Bogen, M. (2018). Help wanted: An examination of hiring algorithms, equity, and bias. Upturn. Retrieved on December 12, 2026, from <https://www.upturn.org/reports/2018/hiring-algorithms/>.

AI ETHICS IN HUMAN RESOURCE MANAGEMENT: A SYSTEMATIC REVIEW

- Sabrina, R., & Zao, T. (2023). Ethical AI in HR: A case study of tech hiring practices. *International Journal of Business Ethics. Semantic Scholar*. <https://www.semanticscholar.org/paper/Ethical-AI-in-HR%3A-A-Case-Study-of-Tech-Hiring-Rao-Zhao/5731fa498375a263ee795e29a4cab9cacc9b3c0e>.
- Selbst, A. D., Boyd, D., Friedler, S. A., Venkatasubramanian, S., and Vertesi, J. (2019). Fairness and abstraction in sociotechnical systems. *In Proceedings of the ACM Conference on Fairness, Accountability, and Transparency*. <https://doi.org/10.1145/3287560.3287598>
- Singh, R., Joshi, A., Dissanayake, H., Nainanayake, D., & Kumar, V. (2025). Harnessing Artificial Intelligence and Human Resource Management for Circular Economy and Sustainability: A Conceptual Integration. *Sustainability*, 17(15), 7054. <https://doi.org/10.3390/su17157054>.
- 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.
- UNESCO. (2021). Recommendation on the ethics of artificial intelligence. Retrieved on November 12, 2025, from <https://unesdoc.unesco.org/ark:/48223/pf0000381137>.
- Varma, A., Dawkins, C., & Chaudhuri, K. (2023). Artificial intelligence and people management: A critical assessment through the ethical lens. *Journal of Business Research*, 157, 113–128.
- Zhang, M.M., Cooke, F.L., Ahlstrom, D. and McNeil, N. (2025), The Rise of Algorithmic Management and Implications for Work and Organisations. *New Technology, Work and Employment*, 40: 659-671. <https://doi.org/10.1111/ntwe.12343>.
- Zhao, H., Yuan, B., Xie, H., & Liao, Y. (2026). Understanding work platformization and algorithmic control from labour process theory. *The Service Industries Journal*, 46(3–4), 332–368. <https://doi.org/10.1080/02642069.2025.2503262>.
- Zhu, C. (2025). The role of AI tools in optimising human resource management practices. *Advances in Curriculum Design & Education*, 1(2), 132. <https://doi.org/10.63808/acde.v1i2.132>.