



DOI: <https://doi.org/10.71317/kjard.2.6.2026.433>

The Kashmir Journal of Academic Research and Development

Journal homepage: <https://rjsaonline.org/index.php/KJARD>



When Algorithms Decide: Fairness, Literacy, and Trust in AI-Powered HR Systems

Saif ur Rehman

Nazeer Hussain University, Karachi, Pakistan

Email: saifur.rehman@nhu.edu.pk

Osama Ahmed

Nazeer Hussain University, Karachi, Pakistan

Email: osama.ahmed@nhu.edu.pk

Farrukh Zafar

Nazeer Hussain University, Karachi, Pakistan

Email: farrukh.zafar@nhu.edu.pk

Shah Salman

Nazeer Hussain University, Karachi, Pakistan

Email: shah.salman@nhu.edu.pk

Einas Azhar Siddiqui

Nazeer Hussain University, Karachi, Pakistan

Email: einas.azhar@nhu.edu.pk

ARTICLE INFO

Received:

May 03, 2026

Revised:

May 11, 2026

Accepted:

June 01, 2026

Available Online:

June 18, 2026

Keywords:

Artificial Intelligence, AI-Driven HR Decision-Making, Perceived Fairness, AI Literacy, Employee Trust, PLS-SEM, Moderated Mediation

Corresponding Author:

Farrukh Zafar

Email:

farrukh.zafar@nhu.edu.pk

ABSTRACT

The rapid diffusion of artificial intelligence (AI) into human resource management (HRM) has shifted decision authority in recruitment, performance evaluation, and promotion from human managers to algorithmic systems, raising urgent questions about how employees come to trust organizations that rely on such systems. Grounded in organizational justice theory and the technology acceptance literature, this study examines the effect of AI-driven HR decision-making on employee trust in the organization, testing the mediating role of perceived fairness and the moderating role of AI literacy on the first-stage path (AI-driven HR decision-making → perceived fairness). Survey data were collected from 300 employees in organizations that have implemented AI-supported HR systems and analyzed using Partial Least Squares Structural Equation Modeling (PLS-SEM) in SmartPLS. The measurement model demonstrated strong reliability and validity (Cronbach's alpha, composite reliability, and AVE all above recommended thresholds; discriminant validity confirmed via the Fornell-Larcker criterion and HTMT ratio). The structural model explained 67.4% of the variance in perceived fairness and 67.6% of the variance in employee trust. Results showed that AI-driven HR decision-making significantly and positively predicted both perceived fairness ($\beta = 0.621$, $p < 0.001$) and employee trust ($\beta = 0.377$, $p < 0.001$), and that perceived fairness significantly mediated this relationship (indirect effect = 0.314, $p < 0.001$). AI literacy significantly moderated the AI-driven HR decision-making-perceived fairness path ($\beta = 0.182$, $p < 0.001$), and moderated mediation was confirmed (index = 0.092, $p < 0.001$), indicating that the indirect effect of AI-driven decision-making on trust through fairness strengthens as employees' AI literacy increases. Blindfolding analysis confirmed adequate predictive relevance for both endogenous constructs ($Q^2 > 0$). The findings suggest that organizations cannot assume AI-driven HR decisions will be trusted by default; rather, trust is built through perceived procedural fairness, and this fairness perception is more readily formed among employees who understand how AI systems work. Theoretical and managerial implications for AI governance, transparency, and workforce AI-literacy training are discussed.

Introduction

Background of the Study

Artificial intelligence (AI) has rapidly transformed the way organizations manage human resources and make employment-related decisions. AI-enabled and algorithmic systems are increasingly being incorporated into recruitment, employee screening, performance evaluation, workforce planning, promotion, training, and other HR functions. These technologies allow organizations to process large amounts of employee information, identify patterns, standardize procedures, and support decisions more efficiently than many traditional approaches. Recent reviews suggest that AI can enhance the efficiency, consistency, objectivity, and accessibility of HR processes, but its organizational consequences depend heavily on how such systems are designed, implemented, governed, and perceived by employees (Naoum et al., 2026; Wagan & Sidra, 2025).

The increasing involvement of AI in HR decision-making is particularly significant because HR decisions directly affect employees' careers, rewards, development opportunities, performance evaluations, and relationships with their organizations. Unlike the use of AI for routine administrative activities, AI-driven HR decision-making involves decisions that may have substantial personal and professional consequences. Consequently, employees may evaluate not only the outcomes produced by AI systems but also whether the process through which these decisions are reached is understandable, unbiased, consistent, and fair. Recent research indicates that employee and applicant reactions to algorithmic HR decisions can be unfavorable, particularly with respect to fairness, justice, trust, organizational attractiveness, and related attitudes (Moritz et al., 2026). Their meta-analysis covering 365 effect sizes from 73 samples and 53 studies found that algorithmic HR decision-making was negatively associated with system-related reactions such as justice, fairness, trust, and trustworthiness.

Therefore, perceived fairness represents a particularly important psychological mechanism for understanding employees' responses to AI-driven HR decisions. Although algorithms may theoretically increase consistency by applying standardized decision rules, employees do not necessarily perceive algorithmic decisions as fair simply because they are data-driven. Perceived fairness can depend on the information used by the system, the procedures through which the algorithm operates, the outcome received by an individual, the availability of explanations, and employees' opportunities to question or challenge decisions. Earlier research on algorithmic decision-making has similarly demonstrated that fairness perceptions are shaped by algorithmic outcomes, development and deployment procedures, and individual differences (Wang et al., 2020).

This issue has become increasingly important because AI in HRM presents a potential paradox. On the one hand, AI systems may reduce some forms of subjective human judgment, improve standardization, increase consistency, and potentially contribute to more objective decision-making. On the other hand, algorithms can reproduce biases embedded in historical data, operate through processes that employees may not understand, and create uncertainty regarding responsibility and accountability. Naoum et al. (2026), for example, concluded from a systematic review of 43 peer-reviewed studies that AI can promote standardization and objectivity in HR processes while simultaneously creating risks associated with embedded bias and weakened accountability. Similarly, Kekez et al. (2025) highlighted continuing conceptual and empirical concerns surrounding bias and discrimination in AI-enabled HRM.

These perceptions of fairness may subsequently influence employee trust in the organization. Trust is particularly important in employment relationships because employees must often accept decisions made by organizational authorities without possessing complete information regarding how those decisions were reached. When employees perceive organizational decision-making processes as fair, they may be more willing to believe that the organization acts responsibly, competently, and with appropriate consideration for employees. Conversely, when AI-supported HR decisions are perceived as opaque, biased, impersonal, or difficult to challenge, employees may question not only the technology but also the organization responsible for adopting and using it.

Empirical evidence increasingly supports this fairness-trust mechanism. Zhou et al. (2021), in an experimental HR recruitment context, found that higher levels of algorithmic fairness were associated with stronger perceptions of fairness and greater user trust in algorithmic decision-making. More directly, Moritz and Schmidt (2024) found that algorithmic management was negatively associated with justice and trust perceptions and, importantly, demonstrated that justice perceptions mediated the relationship between algorithmic management and trust. These findings indicate that employees may evaluate AI not simply as a technological system but as an extension of organizational decision-making, making perceived fairness a potential mechanism through which AI-driven HR decisions influence trust.

However, employees are unlikely to react uniformly to AI-driven HR decision-making. One potentially important individual-level boundary condition is AI literacy. Long and Magerko (2020) conceptualized AI literacy as a collection of competencies that

enables individuals to critically evaluate AI technologies, communicate and collaborate effectively with AI, and use AI appropriately. Thus, AI literacy involves more than knowing how to operate an AI application. It also concerns understanding AI capabilities and limitations and being able to critically interpret AI-supported outputs.

In an HR context, employees with greater AI literacy may be better equipped to understand how AI systems use information, why AI-supported decisions may differ from purely human judgments, what limitations such systems possess, and how AI-generated recommendations should be interpreted. As a result, AI literacy may influence how employees translate exposure to AI-driven HR decision-making into judgments about fairness. However, higher AI literacy should not automatically be assumed to produce more favorable reactions. Greater knowledge may enable employees to recognize both the capabilities and the limitations, biases, and risks of AI systems. This makes AI literacy theoretically valuable as a moderating variable rather than simply assuming that all employees respond to AI-based HR decisions in the same way.

Accordingly, the present study develops an integrated framework in which AI-driven HR decision-making influences employee trust in the organization, both directly and indirectly through perceived fairness, while AI literacy is examined as a boundary condition affecting the relationship between AI-driven HR decision-making and perceived fairness. This framework shifts attention from the technical capabilities of AI toward employees' psychological interpretation of AI-supported organizational decisions.

Problem Statement

Organizations are increasingly adopting AI-driven technologies in HR decision-making because of their potential to improve efficiency, consistency, scalability, and data-driven decisions. However, HR decisions directly affect employees' careers, compensation, performance, promotion, and job security, making employees' perceptions of these technologies particularly important. Moritz et al. (2026) found that algorithmic HR decision-making can negatively affect employees' perceptions of fairness, justice, and trust, suggesting that technological efficiency alone may not ensure positive employee responses.

A key concern is that an AI system considered objective by an organization may still be perceived as unfair by employees because of limited transparency, lack of human consideration, or difficulty in challenging algorithmic decisions. Fairness perceptions depend on both technological characteristics and employees' subjective evaluations (Starke et al., 2021; Wang et al., 2020). Such perceptions may further influence employees' trust in the organization implementing these systems, with evidence suggesting that perceived justice can explain the relationship between algorithmic management and trust (Moritz & Schmidt, 2024).

Furthermore, employees differ in their ability to understand and evaluate AI technologies. AI literacy enables individuals to understand and critically assess AI systems (Long & Magerko, 2020), yet its role in shaping fairness perceptions of AI-driven HR decisions remains underexplored. Therefore, this study addresses the limited understanding of how AI-driven HR decision-making affects employee trust through perceived fairness and whether AI literacy moderates the relationship between AI-driven HR decision-making and perceived fairness.

Research Gap

Despite growing research on AI and algorithmic decision-making in HRM, important gaps remain. Existing studies have primarily examined fairness, justice, and trust as separate employee reactions to algorithmic HR decisions, with limited attention to the mechanisms explaining how AI-driven HR decision-making influences broader employee trust in the organization (Moritz et al., 2026). Although perceived fairness is considered important in evaluating algorithmic decisions (Starke et al., 2021; Wang et al., 2020), its mediating role between AI-driven HR decision-making and employee trust remains insufficiently explored (Moritz & Schmidt, 2024).

Furthermore, AI can enhance consistency and objectivity while simultaneously creating concerns regarding bias, transparency, and accountability (Naoum et al., 2026). Employees may therefore respond differently depending on their ability to understand and evaluate AI systems. Although AI literacy enables individuals to critically assess AI technologies (Long & Magerko, 2020), limited research has investigated AI literacy as a moderator between AI-driven HR decision-making and perceived fairness. Recent studies also call for more employee-centered empirical research on AI-HRM and its psychological consequences (Naoum et al., 2026; Wagan & Sidra, 2025). Therefore, this study addresses these gaps by examining an integrated moderated mediation model, with perceived fairness as a mediator and AI literacy as a moderator in the relationship between AI-driven HR decision-making and employee trust.

Purpose of the Study

The primary purpose of this study is to examine the influence of AI-driven HR decision-making on employee trust in the organization by investigating the psychological mechanism and individual-level boundary condition underlying this relationship.

Specifically, the study seeks to determine whether employees' exposure to AI-driven HR decision-making influences their perceived fairness and subsequent trust in the organization. The study proposes that perceived fairness acts as a mediating mechanism, explaining how AI-driven HR decision-making translates into employee trust. When employees perceive AI-supported HR decisions as fair, consistent, and appropriate, they may be more likely to maintain or develop trust in the organization responsible for those decisions. Conversely, unfavorable fairness perceptions may weaken such trust.

The study further aims to examine AI literacy as a moderating variable. Employees differ in their knowledge and understanding of AI technologies, and these differences may influence how they interpret AI-supported HR decisions. By examining AI literacy as a moderator of the relationship between AI-driven HR decision-making and perceived fairness, the research recognizes that employees are active evaluators of AI rather than passive recipients of algorithmic decisions.

Thus, the study develops a comprehensive model integrating AI-driven HR decision-making, perceived fairness, employee trust, and AI literacy to provide a deeper understanding of employee responses to AI-enabled HR practices. The findings are expected to provide theoretical implications for emerging AI-HRM research and practical guidance for organizations seeking to introduce AI into HR decision-making without undermining perceptions of fairness and organizational trust.

Research Questions

RQ1: Does AI-driven HR decision-making significantly affect employee trust in the organization?

RQ2: Does AI-driven HR decision-making significantly affect employees' perceived fairness?

RQ3: Does perceived fairness significantly affect employee trust in the organization?

RQ4: Does perceived fairness mediate the relationship between AI-driven HR decision-making and employee trust in the organization?

RQ5: Does AI literacy moderate the relationship between AI-driven HR decision-making and perceived fairness?

Research Objectives

RO1: To examine the effect of AI-driven HR decision-making on employee trust in the organization.

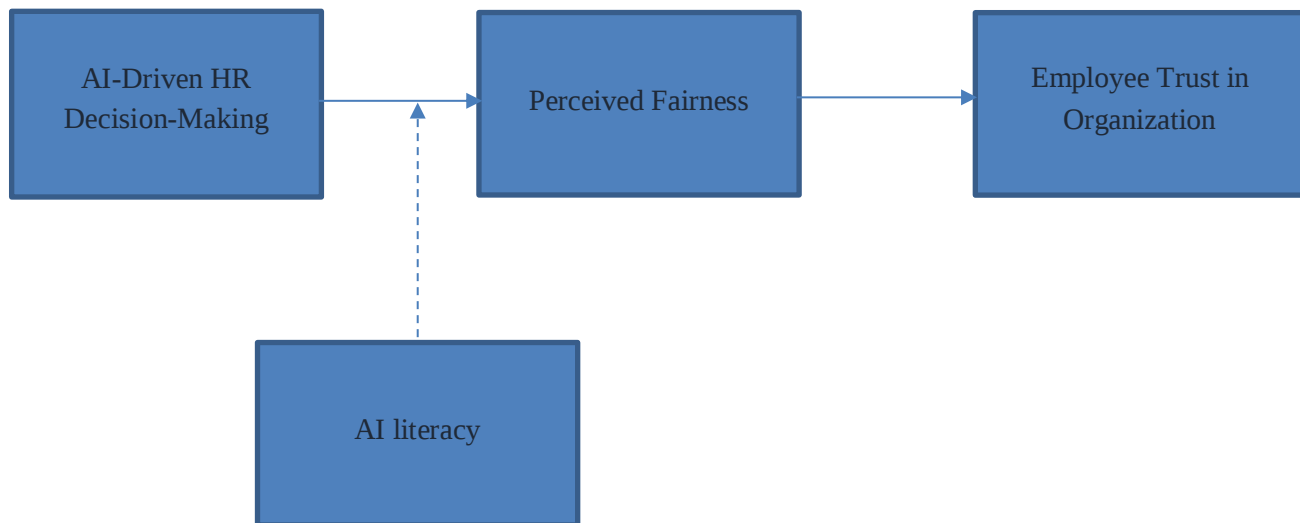
RO2: To examine the effect of AI-driven HR decision-making on employees' perceived fairness.

RO3: To examine the effect of perceived fairness on employee trust in the organization.

RO4: To examine the mediating role of perceived fairness in the relationship between AI-driven HR decision-making and employee trust in the organization.

RO5: To examine the moderating role of AI literacy in the relationship between AI-driven HR decision-making and perceived fairness.

Conceptual Framework



Literature Review

AI-Driven HR Decision-Making and Employee Trust

Artificial intelligence has become increasingly embedded in human resource management, supporting organizational decisions related to recruitment, selection, performance evaluation, promotion, workforce planning, and employee development. AI-driven HR decision-making enables organizations to analyze large volumes of employee data, standardize decision criteria, and potentially reduce inconsistencies associated with purely human judgment. Recent research suggests that AI can improve efficiency and objectivity in HRM, although its effects depend considerably on how employees perceive and experience algorithmic systems (Naoum et al., 2026). In particular, the growing involvement of AI in consequential employment decisions has raised questions about whether employees are willing to trust organizations that delegate or partially delegate HR decisions to algorithms.

Trust is critical in AI-enabled HRM because employees generally lack complete information about how algorithms reach recommendations or decisions. When employees believe that AI-driven HR systems are reliable, appropriately implemented, and used responsibly, they may develop more favorable perceptions of the organization employing them. In contrast, opaque algorithms, limited human involvement, perceived bias, and uncertainty regarding accountability may reduce trust. Strohmeier et al. (2026) identify fairness, accountability, and transparency as fundamental ethical principles for appropriate algorithmic decision-making in HRM, indicating that employee acceptance depends on more than the technical performance of the system.

Empirical findings indicate that the relationship between algorithmic management and trust can be complex. Moritz and Schmidt (2024) found that algorithmic management was negatively associated with employees' trust perceptions in their experimental study, suggesting that employees may perceive automated management as less trustworthy under certain circumstances. Similarly, research on AI-assisted HRM highlights trust as a central challenge when AI becomes involved in performance evaluation and other employee-related decisions (Nguyen & Connolly, 2025). These findings demonstrate that AI-driven decision-making can meaningfully shape how employees evaluate the organization itself. Nevertheless, when AI systems are implemented transparently, consistently, and responsibly, they may also provide employees with greater confidence that HR decisions are evidence-based rather than dependent entirely on subjective managerial judgment. Accordingly, AI-driven HR decision-making is expected to have a significant relationship with employee trust.

H1: AI-driven HR decision-making has a significant positive effect on employee trust in the organization.

AI-Driven HR Decision-Making and Perceived Fairness

Perceived fairness represents employees' subjective assessment of whether organizational decisions, procedures, and outcomes are reasonable, unbiased, consistent, and appropriate. Fairness is particularly relevant to AI-driven HR decision-making because algorithms increasingly participate in decisions that directly affect employees' careers. Employees may evaluate whether the information used by an AI system is relevant, whether similar employees are treated consistently, whether the decision process is transparent, and whether opportunities exist to challenge potentially incorrect decisions.

AI-driven HR decision-making has the potential to improve fairness by applying standardized criteria to employees and reducing some forms of subjective human judgment. Qin et al. (2023), for example, conducted a field experiment comparing AI and human managers in employee performance evaluation and found that employees perceived AI as both fairer and more accurate than the average human manager. Their findings indicate that employees do not necessarily view AI involvement negatively and may recognize the capacity of AI to provide standardized and data-intensive evaluations.

However, the use of AI does not automatically guarantee perceived fairness. AI systems may reproduce biases contained in historical organizational data, and employees may perceive decisions as unfair when algorithms are opaque or when employees cannot understand, question, or appeal their outcomes. Naoum et al. (2026) describe this as a dual impact of AI in HRM: AI may enhance standardization and objectivity while simultaneously introducing concerns regarding embedded bias and accountability. Strohmeier et al. (2026) similarly identify fairness as one of the essential ethical principles required for appropriate algorithmic HR decision-making.

The way AI systems are designed and implemented is therefore critical to employees' fairness perceptions. When AI-supported HR decisions are based on relevant information, follow consistent criteria, and reduce arbitrary managerial judgments, employees may consider them fairer. In contrast, poorly designed or opaque systems can produce unfavorable fairness judgments. This indicates that employees' perceptions of AI-driven HR decision-making are directly connected to how fair they consider organizational HR processes to be.

H2: AI-driven HR decision-making has a significant positive effect on employees' perceived fairness.

Perceived Fairness and Employee Trust in the Organization

Employee trust reflects an individual's willingness to rely on an organization based on expectations that it will act competently, consistently, ethically, and with appropriate consideration for employees. Within HRM, fairness is an important foundation for developing such trust because employees continuously evaluate organizational intentions through HR practices and decisions. Fair treatment can communicate that the organization values employees and follows acceptable standards, whereas unfair treatment may signal opportunistic or irresponsible organizational behavior.

This relationship can be explained through organizational justice principles. When employees believe that decision procedures are consistent and unbiased, outcomes are appropriate, explanations are adequate, and they are treated respectfully, they are more likely to regard organizational authorities as trustworthy. In an AI-enabled environment, fairness becomes particularly important because employees may have limited direct interaction with the mechanism responsible for a decision. Consequently, their assessment of whether the process was fair can serve as an important basis for determining whether the organization behind that process deserves their trust.

Moritz and Schmidt (2024) provide empirical support for this relationship. Their experimental research found that algorithmic management was associated with both justice and trust perceptions and demonstrated that justice perceptions explained part of the effect of algorithmic management on trust. Similarly, Nguyen and Connolly (2025) emphasize organizational justice and organizational trust as important components of employee trust formation in AI-assisted HR performance evaluation.

Recent research further shows that fairness remains highly relevant when organizations implement AI-based HR practices. Irvan et al. (2026), examining employees exposed to AI-enabled HR decision processes, identified algorithmic fairness and trust in management as central variables for understanding employee responses to AI-based HR systems. Therefore, employees who perceive AI-supported HR practices as fair should be more willing to trust the organization responsible for implementing and governing those technologies.

H3: Perceived fairness has a significant positive effect on employee trust in the organization.

Mediating Role of Perceived Fairness

Although AI-driven HR decision-making may directly influence employee trust, the relationship is unlikely to occur exclusively through direct technological effects. Employees interpret organizational technologies and form psychological evaluations before developing broader attitudes toward their organization. Perceived fairness provides a particularly relevant mechanism because AI-driven HR decisions communicate how organizational resources, opportunities, rewards, evaluations, and employment outcomes are allocated.

When employees encounter AI-driven HR decision-making, they may initially evaluate whether the process is fair rather than immediately deciding whether the organization itself is trustworthy. If the AI system applies consistent criteria, provides reliable assessments, reduces arbitrary managerial judgments, and treats employees equitably, employees may develop stronger fairness perceptions. These fairness judgments can subsequently increase confidence in the organization responsible for selecting, implementing, and monitoring the AI system. Conversely, an opaque, biased, or unchallengeable AI decision may create perceptions of unfairness that subsequently weaken organizational trust.

This argument is supported by research showing that fairness is central to employees' reactions to algorithmic management. Moritz and Schmidt (2024) found that algorithmic management was negatively related to justice and trust perceptions and, importantly, that justice perceptions partially mediated the relationship between algorithmic management and trust. Their findings provide direct evidence that employees' fairness-related judgments constitute an explanatory pathway through which algorithmic management affects trust.

The mediating argument is also consistent with evidence that employees can perceive AI as fairer than human decision-makers under certain circumstances. Qin et al. (2023) found that AI-based employee performance evaluations were perceived as fairer and more accurate than evaluations conducted by the average human manager. Therefore, when AI-driven HR decision-making enhances employees' perceptions of consistency and fairness, these perceptions may translate into greater trust in the organization. Perceived fairness thus represents a psychological bridge connecting employees' experiences of AI-driven HR decisions with their broader assessment of organizational trustworthiness.

H4: Perceived fairness significantly mediates the relationship between AI-driven HR decision-making and employee trust in the organization.

Moderating Role of AI Literacy

Employees are unlikely to interpret AI-driven HR decisions uniformly. An important source of this variation is AI literacy, which refers to an individual's knowledge and competencies for understanding, evaluating, interacting with, and critically assessing artificial intelligence. AI literacy is particularly important in contemporary workplaces because employees increasingly encounter AI systems without necessarily possessing technical expertise regarding how algorithms generate predictions, recommendations, or decisions.

Long and Magerko (2020) conceptualize AI literacy as a set of competencies that enables individuals to critically evaluate AI technologies, communicate and collaborate effectively with AI, and use AI appropriately. Employees with greater AI literacy may therefore be better positioned to understand both the capabilities and limitations of AI-driven HR systems. Rather than interpreting every algorithmic outcome as an unquestionable objective judgment or an unexplained technological decision, AI-literate employees may be more capable of evaluating the information used, limitations of algorithmic predictions, and role of human oversight.

Emerging empirical research demonstrates the importance of AI literacy in workplace decision processes. Zhao and Shi (2026), studying 435 employees, found that AI literacy was positively associated with employee-AI collaboration and decision-making quality, while AI trust functioned as an important boundary condition. These findings suggest that employees' knowledge and understanding of AI can influence how they engage with and evaluate AI-supported decision-making.

More directly within HRM, research on recruiting recommender systems suggests that AI literacy influences how HR users interpret explainable AI. Kalff and Simbeck (2025), in an experiment involving 410 HR managers, found that AI literacy affected responses to AI explanations. Explainable AI features improved subjective helpfulness and trust particularly among individuals with moderate or high AI literacy, although greater explanation did not necessarily guarantee greater objective understanding. This indicates that the same AI system or explanation can produce different reactions depending on the user's AI literacy.

These findings provide a basis for considering AI literacy as a moderator rather than merely an independent predictor. Employees with high AI literacy may better understand why AI is used in HR decisions, how algorithmic recommendations are generated, and where human oversight remains necessary. This understanding can reduce uncertainty and enable more informed fairness evaluations. When AI-driven HR systems operate consistently and responsibly, employees with higher AI literacy may be better able to recognize these characteristics and consequently perceive the decision process as fair. Employees with lower AI literacy, in contrast, may experience greater uncertainty regarding how the system operates and may find it more difficult to assess the legitimacy of its decisions.

At the same time, AI literacy should not be interpreted as unconditional acceptance of AI. More knowledgeable employees may also be better able to recognize algorithmic bias, data limitations, and inappropriate applications. Thus, AI literacy influences the strength of the relationship between AI-driven HR decision-making and perceived fairness by changing employees' capacity to interpret and critically evaluate AI-supported decisions. Based on the expectation that appropriate AI-driven HR decision-making will be more positively understood by employees with stronger AI literacy, the following hypothesis is proposed:

H5: AI literacy significantly and positively moderates the relationship between AI-driven HR decision-making and perceived fairness, such that the positive relationship is stronger when employees' AI literacy is high.

Methodology

Research Design

This study adopts a quantitative, cross-sectional, and explanatory research design to examine the relationships among AI-driven HR decision-making, perceived fairness, employee trust in the organization, and AI literacy. A quantitative approach is appropriate because the study aims to empirically test theoretically developed hypotheses and determine the strength and significance of relationships among the proposed constructs. Data will be collected from respondents at a single point in time through a structured questionnaire. The proposed model incorporates both mediation and moderation effects; therefore, Partial Least Squares Structural Equation Modeling (PLS-SEM) will be employed to test the hypothesized relationships. PLS-SEM is particularly appropriate for examining complex models involving direct, indirect, and interaction effects and for prediction-oriented research (Hair et al., 2022).

Population and Sampling

The target population of this study consists of employees working in organizations where artificial intelligence or AI-enabled systems are used to support HR-related activities and decisions, such as recruitment and selection, performance evaluation, promotion, training, workforce planning, or employee management. Employees with exposure to AI-supported HR practices are considered appropriate respondents because they can meaningfully evaluate the fairness of such systems and their influence on organizational trust.

A purposive sampling technique will be employed to ensure that respondents have relevant experience or awareness of AI-enabled HR practices. A screening question will be included at the beginning of the questionnaire to confirm respondents' exposure to AI-supported HR processes. The required sample size will be determined considering the complexity of the structural model and recommendations for PLS-SEM. Following Hair et al. (2022), an adequate sample will be targeted to ensure sufficient statistical power for estimating direct, mediation, and moderation effects.

Data Collection Procedure

Data will be collected using a structured self-administered questionnaire distributed both electronically and, where appropriate, in printed form. Before the main data collection, respondents will be informed about the academic purpose of the study, voluntary participation, confidentiality, and anonymity of their responses. A brief explanation of AI-driven HR decision-making will be provided to ensure that respondents understand the context of the study.

The questionnaire will initially contain screening and demographic questions, followed by measurement items for AI-driven HR decision-making, perceived fairness, employee trust, and AI literacy. Respondents will be requested to indicate their level of agreement with each statement using a five-point Likert scale ranging from 1 = strongly disagree to 5 = strongly agree. Before full-scale data collection, the questionnaire may be pilot-tested with a small group of respondents to evaluate clarity, wording, reliability, and comprehensibility.

Measurement Instrument

The study will use a structured questionnaire consisting primarily of measurement items adapted from previously established scales. Minor modifications may be made to the wording of items to ensure their relevance to the context of AI-driven HR decision-making while retaining the conceptual meaning of the original measures. All constructs will be operationalized as reflective constructs.

Construct	Source
AI-Driven HR Decision-Making (AI-HRDM)	Adapted from Brougham and Haar (2018) and Köchling and Wehner (2020)
Perceived Fairness (PF)	Adapted from Colquitt (2001) and Wang et al. (2020)
Employee Trust in the Organization (ET)	Adapted from Nyhan and Marlowe (1997) and Robinson (1996)
AI Literacy (AIL)	Adapted from Long and Magerko (2020) and validated AI literacy measures developed by Ng et al. (2021)

The questionnaire items will be adapted to explicitly reflect the AI-HRM context. For instance, perceived fairness items will assess whether employees consider AI-supported HR decisions consistent, unbiased, appropriate, and fair, while employee trust items will capture employees' confidence in the organization responsible for adopting and implementing AI-supported HR decisions. AI literacy items will assess employees' knowledge, understanding, and ability to critically evaluate AI technologies.

Data Analysis Technique

Data will be analyzed using SmartPLS and the PLS-SEM technique. Before conducting the main analysis, the collected data will be screened for missing values, incomplete responses, response patterns, and potential data-entry errors. Descriptive statistics will be used to summarize respondents' demographic characteristics. PLS-SEM analysis will subsequently be conducted in two stages: (1) measurement model assessment and (2) structural model assessment, following the recommendations of Hair et al. (2022).

Measurement Model Assessment

The measurement model will be assessed to establish the reliability and validity of the constructs before testing the structural relationships. First, indicator reliability will be evaluated using outer loadings, with values of approximately 0.708 or above considered desirable. Internal consistency reliability will be assessed through Cronbach's alpha, rho_A, and composite reliability (CR). Values between 0.70 and 0.95 will generally indicate satisfactory internal consistency.

Convergent validity will be examined using the Average Variance Extracted (AVE). An AVE value of 0.50 or higher indicates that a construct explains at least 50% of the variance in its indicators (Hair et al., 2022). Discriminant validity will primarily be assessed through the Heterotrait-Monotrait ratio (HTMT). HTMT values below 0.85, or below 0.90 under a more liberal criterion, will indicate adequate discriminant validity (Henseler et al., 2015). The Fornell-Larcker criterion may also be reported as supplementary evidence, where the square root of each construct's AVE should exceed its correlations with other constructs.

Structural Model Assessment

After establishing satisfactory reliability and validity, the structural model will be evaluated to test the proposed hypotheses. First, potential multicollinearity among predictor constructs will be examined using the Variance Inflation Factor (VIF). The structural relationships will then be evaluated through standardized path coefficients (β). The statistical significance of the hypothesized relationships will be determined using the bootstrapping procedure, preferably with 5,000 bootstrap subsamples. The t-values, p-values, and confidence intervals will be examined to determine whether the proposed hypotheses are supported. The model's explanatory power will be assessed using the coefficient of determination (R^2) for endogenous constructs, particularly perceived fairness and employee trust. The effect size (f^2) will be examined to determine the relative contribution of individual predictors, while Q^2 will be considered to assess the model's predictive relevance.

For the mediation analysis (H4), the specific indirect effect of:

AI-Driven HR Decision-Making → Perceived Fairness → Employee Trust

will be tested through bootstrapping. A statistically significant indirect effect will provide evidence of mediation. The direct and indirect effects will also be considered to determine whether the mediation is partial or full.

For the moderation analysis (H5), an interaction term between AI-driven HR decision-making and AI literacy (AI-HRDM × AIL) will be created to predict perceived fairness. A statistically significant interaction effect will provide evidence of moderation. Simple slope analysis may subsequently be conducted to illustrate how the relationship between AI-driven HR decision-making and perceived fairness changes at low and high levels of AI literacy.

Findings

Measurement Analysis

The measurement model was assessed to establish the reliability, convergent validity, and discriminant validity of the study constructs: AI Literacy (AIL), AI-Driven HR Decision-Making (AI-HRDM), Employee Trust in Organization (ETO), and Perceived Fairness (PF). The assessment included Cronbach's alpha, rho_A, composite reliability (CR), Average Variance Extracted (AVE), Fornell-Larcker criterion, HTMT, and model fit indices. The results demonstrate satisfactory measurement quality and provide sufficient support for proceeding with structural model assessment.

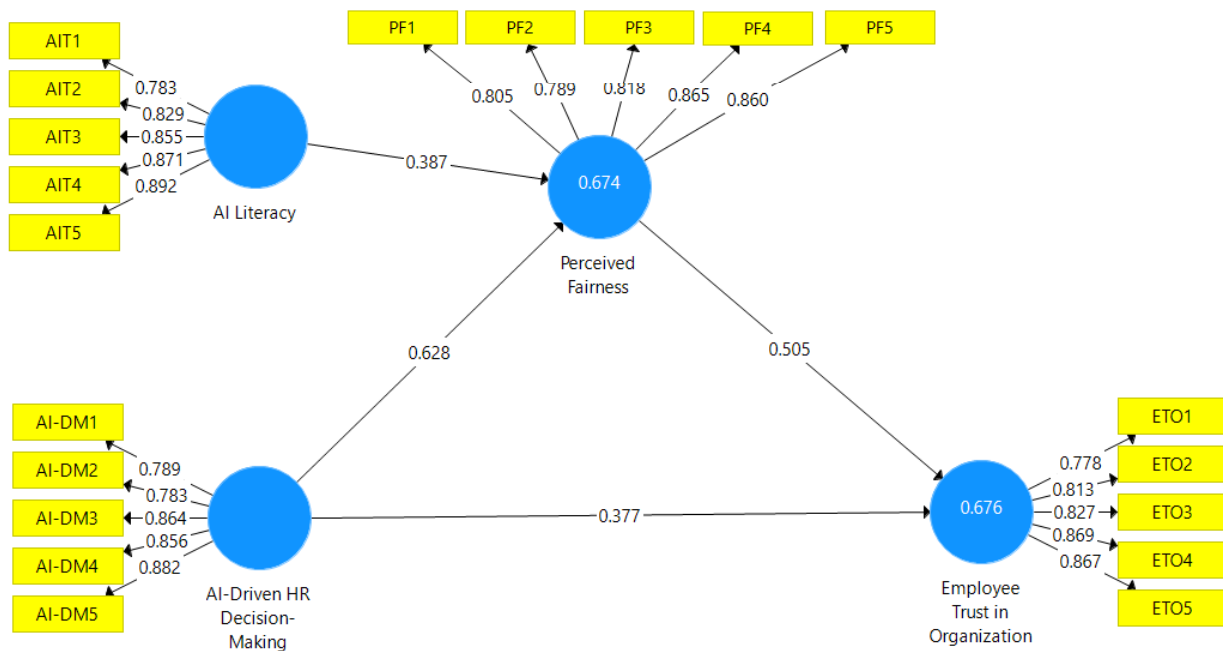


Figure 1. Measurement Model

Table 1. Construct Reliability and Validity

Construct	Cronbach's Alpha	rho_A	Composite Reliability	AVE
AI Literacy	0.901	0.912	0.927	0.717
AI-Driven HR Decision-Making	0.892	0.896	0.920	0.699
Employee Trust in Organization	0.888	0.893	0.918	0.691

Perceived Fairness	0.885	0.889	0.916	0.686
--------------------	-------	-------	-------	-------

The results in Table 1 demonstrate satisfactory internal consistency reliability, as Cronbach's alpha, rho_A, and composite reliability values for all constructs exceed the recommended threshold of 0.70. Composite reliability values range from 0.916 to 0.927, confirming strong construct reliability. Furthermore, AVE values range from 0.686 to 0.717, exceeding the recommended threshold of 0.50. Therefore, the constructs demonstrate satisfactory convergent validity (Hair et al., 2022).

Table 2. Fornell-Larcker Criterion

Construct	AIL	AI-HRDM	ETO	PF
AIL	0.847			
AI-HRDM	0.266	0.836		
ETO	0.375	0.746	0.832	
PF	0.554	0.731	0.781	0.828

The Fornell-Larcker criterion was used to examine discriminant validity. As shown in Table 2, the square root of AVE for each construct, presented diagonally, is greater than its corresponding inter-construct correlations. For example, the square root of AVE for AI Literacy is 0.847, AI-HRDM is 0.836, Employee Trust is 0.832, and Perceived Fairness is 0.828. These findings provide evidence of satisfactory discriminant validity.

Table 3. Heterotrait-Monotrait Ratio (HTMT)

Construct	AIL	AI-HRDM	ETO	PF
AIL	—			
AI-HRDM	0.294	—		
ETO	0.411	0.835	—	
PF	0.613	0.820	0.877	—

The HTMT criterion was also examined to further establish discriminant validity. All HTMT values are below the threshold of 0.90, ranging from 0.294 to 0.877. The highest HTMT value of 0.877 was observed between Perceived Fairness and Employee Trust, but it remained below the 0.90 threshold. Therefore, adequate discriminant validity among all study constructs was established (Henseler et al., 2015).

Table 4. Model Fit Summary

Fit Index	Saturated Model	Estimated Model
SRMR	0.046	0.046
d_ULS	0.438	0.439
d_G	0.182	0.182
Chi-Square	309.089	309.225

NFI	0.925	0.925
-----	-------	-------

The model fit results indicate satisfactory overall model fit. The SRMR value of approximately 0.046 for both the saturated and estimated models is substantially below the commonly recommended threshold of 0.08. Moreover, the NFI value of approximately 0.925 indicates an acceptable level of model fit. Overall, the measurement model demonstrates adequate reliability, validity, and model fit and is therefore suitable for structural model assessment.

Structural Analysis

After confirming the reliability and validity of the measurement model, the structural model was evaluated to examine the hypothesized relationships. Path coefficients, t-statistics, p-values, coefficient of determination (R^2), predictive relevance (Q^2), mediation, and moderation effects were examined.

Table 5. R^2 and Predictive Relevance

Endogenous Construct	R Square	R Square Adjusted	Q^2
Employee Trust in Organization	0.676	0.674	0.460
Perceived Fairness	0.674	0.672	0.479

The R^2 value for Employee Trust in Organization is 0.676, indicating that approximately 67.6% of the variance in employee trust is explained by the predictors included in the model. Similarly, the R^2 value for Perceived Fairness is 0.674, demonstrating that approximately 67.4% of its variance is explained by AI-driven HR decision-making, AI literacy, and their interaction. These findings demonstrate substantial explanatory power of the proposed model.

Furthermore, the Q^2 values for Employee Trust (0.460) and Perceived Fairness (0.479) are considerably greater than zero, demonstrating that the structural model possesses satisfactory predictive relevance for both endogenous constructs.

Table 6. Specific Direct and Indirect Pathways

Hyp.	Hypothesis	Beta (β)	T Value	P Value	Result
H1	AI-HRDM $\rightarrow$ Employee Trust	0.377	7.566	<0.001	Accepted
H2	AI-HRDM $\rightarrow$ Perceived Fairness	0.621	19.556	<0.001	Accepted
H3	Perceived Fairness $\rightarrow$ Employee Trust	0.505	9.961	<0.001	Accepted
H4	AI-HRDM $\rightarrow$ Perceived Fairness $\rightarrow$ Employee Trust	0.314	8.720	<0.001	Accepted
H5	AI Literacy $\times$ AI-HRDM $\rightarrow$ Perceived Fairness	0.182	6.644	<0.001	Accepted

The structural model results show that all five proposed hypotheses are supported. H1 indicates that AI-driven HR decision-making has a significant positive effect on employee trust ($\beta = 0.377$, $t = 7.566$, $p < 0.001$). Thus, greater positive perceptions of AI-driven HR decision-making are associated with stronger employee trust in the organization.

H2 is also supported, demonstrating a strong positive relationship between AI-driven HR decision-making and perceived fairness ($\beta = 0.621$, $t = 19.556$, $p < 0.001$). This indicates that employees who perceive AI-driven HR decision-making more positively are considerably more likely to view organizational HR processes as fair.

H3 confirms that perceived fairness significantly and positively influences employee trust ($\beta = 0.505$, $t = 9.961$, $p < 0.001$). Thus, employees' perceptions of fairness represent an important determinant of their trust in the organization.

Mediation Analysis

The indirect relationship between AI-driven HR decision-making and employee trust through perceived fairness was statistically significant ($\beta = 0.314$, $t = 8.720$, $p < 0.001$). Therefore, H₄ is supported. Since the direct effect of AI-driven HR decision-making on employee trust also remains positive and statistically significant ($\beta = 0.377$, $p < 0.001$), the results indicate complementary partial mediation. This means that AI-driven HR decision-making influences employee trust both directly and indirectly by strengthening employees' perceptions of fairness.

Moderation Analysis

The interaction between AI Literacy and AI-Driven HR Decision-Making has a significant positive effect on perceived fairness ($\beta = 0.182$, $t = 6.644$, $p < 0.001$). Therefore, H₅ is supported. The positive interaction coefficient indicates that AI literacy strengthens the positive relationship between AI-driven HR decision-making and perceived fairness. In other words, employees with greater AI literacy are more likely to positively interpret AI-driven HR decision-making in terms of fairness.

The analysis also revealed a significant positive direct effect of AI literacy on perceived fairness ($\beta = 0.399$, $t = 11.509$, $p < 0.001$). Thus, besides its moderating role, greater AI literacy is directly associated with stronger fairness perceptions.

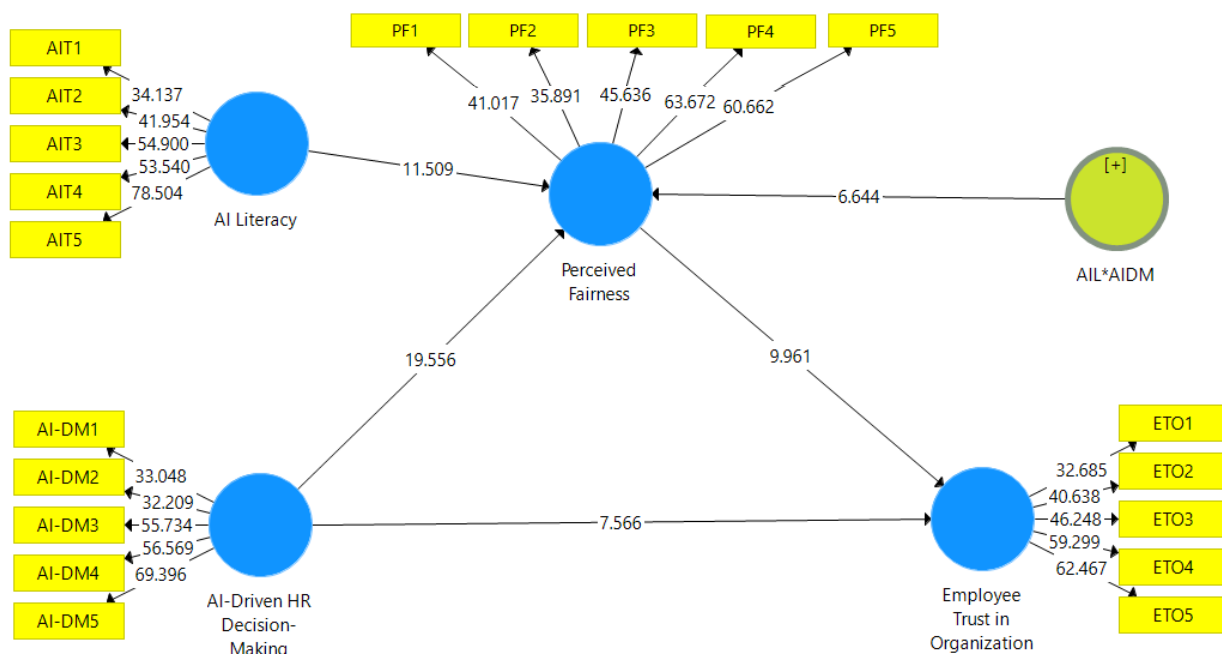


Figure 2. Structural Model with Moderation Effect

Overall, the structural results provide strong empirical support for the proposed research framework. AI-driven HR decision-making directly increases employee trust and perceived fairness, while perceived fairness increases employee trust and partially mediates the AI-HRDM-trust relationship. Furthermore, AI literacy acts as a significant positive moderator, strengthening the effect of AI-driven HR decision-making on perceived fairness.

Conclusion

This study examined the relationships among AI-driven HR decision-making, perceived fairness, employee trust in the organization, and AI literacy, with particular attention to the mediating role of perceived fairness and the moderating role of AI literacy. The findings provide strong empirical support for the proposed research model, as all five hypotheses were accepted. The results demonstrate that AI-driven HR decision-making has a significant positive influence on employee trust in the organization ($\beta = 0.377$, $p < 0.001$) and an even stronger positive influence on employees' perceived fairness ($\beta = 0.621$, $p < 0.001$). These findings indicate that when employees perceive AI-supported HR decisions as systematic, reliable, consistent, and appropriately implemented, they are more likely to consider organizational decision-making processes fair and develop greater

trust in the organization. Thus, the effectiveness of AI in HRM extends beyond operational efficiency and depends considerably on how employees experience and evaluate AI-supported decisions.

The findings further establish perceived fairness as an important psychological mechanism connecting AI-driven HR decision-making with employee trust. Perceived fairness had a significant positive effect on employee trust ($\beta = 0.505$, $p < 0.001$), demonstrating that employees who consider HR decisions fair are more likely to trust the organization responsible for those decisions. Furthermore, the significant indirect effect of AI-driven HR decision-making on employee trust through perceived fairness ($\beta = 0.314$, $p < 0.001$) confirms the mediating role of perceived fairness. Since the direct relationship between AI-driven HR decision-making and employee trust remained significant, the findings indicate partial mediation. This means that AI-driven HR decision-making can directly strengthen organizational trust while also influencing trust indirectly by improving employees' fairness perceptions.

Another important conclusion concerns the role of AI literacy. The significant positive interaction between AI literacy and AI-driven HR decision-making ($\beta = 0.182$, $p < 0.001$) demonstrates that AI literacy strengthens the relationship between AI-driven HR decision-making and perceived fairness. Employees who possess greater knowledge and understanding of AI technologies may be better able to evaluate how AI-supported HR decisions are generated and recognize their potential benefits and limitations. Consequently, they may develop more informed fairness judgments regarding AI-enabled HR practices. The model also demonstrated substantial explanatory and predictive capabilities, explaining 67.6% of the variance in employee trust and 67.4% of the variance in perceived fairness. The positive Q^2 values for employee trust (0.460) and perceived fairness (0.479) further confirm the predictive relevance of the proposed model. Overall, the findings suggest that successful implementation of AI in HRM requires organizations to integrate technological effectiveness with fairness, employee trust, transparency, human involvement, and AI literacy.

Recommendations

Based on the findings, organizations should adopt a human-centered approach to AI-driven HR decision-making rather than viewing AI implementation solely as a technological or efficiency-oriented initiative. AI systems used in recruitment, performance evaluation, promotion, compensation, training, and career development should operate according to clearly defined and consistent criteria. Organizations should regularly evaluate AI-supported decisions for potential bias, discrimination, inaccuracies, and unintended consequences. Employees should also be provided with sufficient information regarding how AI contributes to HR decisions, what types of data are considered, and how final decisions are reached. Increasing transparency and explainability can reduce uncertainty and help employees understand the reasoning behind AI-supported HR outcomes.

Organizations should also maintain appropriate human oversight, particularly for HR decisions that have substantial consequences for employees. AI should support rather than completely replace human judgment in sensitive decisions such as promotion, termination, compensation, and performance evaluation. Employees should have access to mechanisms through which they can request clarification, provide additional information, or challenge an AI-supported decision when they believe that relevant circumstances have not been adequately considered. Such procedures can enhance perceptions of procedural fairness and demonstrate organizational accountability.

Furthermore, organizations should invest systematically in AI literacy development. The significant moderating role of AI literacy suggests that employees' understanding of AI can influence how they evaluate AI-driven HR decisions. Training programs should therefore explain basic AI concepts, how organizational AI systems operate, their capabilities and limitations, potential sources of algorithmic bias, and the role of human oversight. These programs should be accessible to employees at different organizational levels rather than being restricted to HR or technical personnel. Improving AI literacy can enable employees to engage more confidently and critically with AI systems, contributing to stronger fairness perceptions and organizational trust.

Future Research Recommendations

Although the present study provides important evidence regarding employee responses to AI-driven HR decision-making, several opportunities exist for future research. First, future studies could extend the proposed model by incorporating additional technological and organizational variables such as AI transparency, explainability, perceived algorithmic bias, privacy concerns, data security, human oversight, organizational justice, and perceived organizational support. These factors may provide a more comprehensive explanation of why employees accept or resist AI-driven HR practices. Researchers could also investigate

additional mediating mechanisms, such as perceived transparency, perceived control, psychological safety, or perceived organizational justice, to further explain how AI-supported HR decisions influence employee attitudes and behaviors.

Second, future studies should consider longitudinal and experimental research designs. Since the present framework is based on cross-sectional data, longitudinal research could examine how employees' perceptions of AI-driven HR decision-making, fairness, and trust change as they gain greater experience with AI systems. Experimental studies could manipulate characteristics such as AI transparency, human involvement, explainability, or decision consequences to provide stronger evidence regarding causal relationships. Future research could also compare fully automated decisions, human decisions, and hybrid human-AI decisions to identify which decision structure produces the strongest perceptions of fairness and organizational trust.

Third, future studies could test the proposed framework across different industries, organizational contexts, occupations, and countries. Employee reactions to AI-driven HR decision-making may differ between technology-intensive and traditional industries or between public and private organizations. Cultural differences may also influence how employees interpret algorithmic authority, fairness, and organizational trust. Cross-country studies could therefore provide valuable evidence regarding the generalizability of the present findings. Finally, researchers could extend the consequences of AI literacy beyond fairness by examining its influence on AI acceptance, employee engagement, job satisfaction, organizational commitment, innovative behavior, resistance to AI, and employee-AI collaboration.

Managerial Implications

The findings have significant practical implications for HR managers, organizational leaders, and decision-makers responsible for AI implementation. The strong positive relationship between AI-driven HR decision-making and perceived fairness ($\beta = 0.621$) demonstrates that employees' fairness judgments are closely connected to how AI-supported HR systems are designed and implemented. Managers should therefore recognize that technical accuracy alone does not guarantee successful AI adoption. Employees must also perceive the decision-making process as legitimate, consistent, understandable, and fair. Before introducing an AI system into HR processes, organizations should establish clear governance standards concerning data quality, algorithmic bias, transparency, accountability, and human oversight.

The significant relationship between perceived fairness and employee trust ($\beta = 0.505$) has particularly important managerial implications. Employees do not evaluate AI systems independently of the organizations that implement them. If an AI-supported HR decision is perceived as unfair, employees may attribute responsibility to management and consequently lose trust in the organization itself. Managers should therefore establish procedures that allow employees to understand how decisions are made and provide opportunities for review or appeal. HR professionals should also regularly monitor employee perceptions of AI-supported decisions rather than relying exclusively on technical performance indicators. Employee feedback can help identify fairness concerns before they develop into broader resistance or distrust.

The mediating effect of perceived fairness ($\beta = 0.314$) further indicates that fairness represents a key pathway through which organizations can convert AI-driven HR practices into stronger employee trust. Therefore, managers seeking to build trust during digital transformation should prioritize fair AI governance. This includes using relevant and unbiased data, conducting regular algorithmic audits, communicating decision criteria, maintaining human accountability, and ensuring that employees are treated respectfully throughout AI-supported processes. Fairness should consequently be treated as a strategic requirement of AI implementation rather than merely an ethical consideration.

Finally, the significant moderating effect of AI literacy ($\beta = 0.182$) demonstrates that organizations should consider employees' technological capabilities when implementing AI-driven HR systems. Employees with limited understanding of AI may experience uncertainty or difficulty interpreting algorithmic decisions, whereas employees with greater AI literacy may be better positioned to critically evaluate AI-supported processes. Managers should therefore provide continuous AI awareness and literacy programs, practical demonstrations, workshops, and accessible guidance before and during AI implementation. Combining responsible AI governance with employee education can help organizations gain the efficiency benefits of AI while simultaneously strengthening perceived fairness, employee confidence, and long-term organizational trust.

References

- Ahmed, O., Saqib, A., Salman, S., Siddiqui, E. A., & Ali, O. (2025). AI-powered HRM and sustainable organizational performance: The mediating roles of green learning and engagement under environmental uncertainty. *Research Journal for Social Affairs*, 3(6), 1383–1399. <https://doi.org/10.71317/RJSA.003.06.0617>

- Ali, S., Haidery, A., Dawood, J., Kamran, A., & Ahmed, O. (2025, July). Sustainable Banking: Exploring the Interplay of Eco-Friendly Behaviors, Cultural Dynamics, and Green HRM Strategies in Pakistan's Financial Landscape. In International Conference on Management Science and Engineering Management (pp. 726-744). Singapore: Springer Nature Singapore.
- Azher, E., Javed, S., Siddiqui, H. M. A., Zafar, F., & Ahmed, O. (2025). Exploring the impact of digital supply chain integration on the firm's performance with mediation and moderation role of knowledge sharing and environmental turbulence. *Social Science Review Archives*, 3(1), 567-581.
- Brougham, D., & Haar, J. (2018). Smart technology, artificial intelligence, robotics, and algorithms (STARA): Employees' perceptions of our future workplace. *Journal of Management & Organization*, 24(2), 239-257.
- Colquitt, J. A. (2001). On the dimensionality of organizational justice: A construct validation of a measure. *Journal of Applied Psychology*, 86(3), 386-400.
- Haidery, A., Ali, S., Dawood, J., Kamran, A., & Ahmed, O. (2025, July). Impact of Green Life Work Balance on Organization Environment: A Case of Pakistan Insurance Industry. In International Conference on Management Science and Engineering Management (pp. 388-402). Singapore: Springer Nature Singapore.
- Hair, J. F., Hult, G. T. M., Ringle, C. M., & Sarstedt, M. (2022). *A primer on partial least squares structural equation modeling (PLS-SEM)* (3rd ed.). Sage.
- Henseler, J., Ringle, C. M., & Sarstedt, M. (2015). A new criterion for assessing discriminant validity in variance-based structural equation modeling. *Journal of the Academy of Marketing Science*, 43, 115-135.
- Irvan, M., Syahroni, B., & Mahadianto, M. Y. (2026). Transparency, algorithmic fairness, and employee performance in AI-based HR: The moderating role of trust in management. *Indonesian Journal Economic Review*, 6(1), 101-115.
- Javed, S., Azher, E., Jabbar, B., & Zafar, F. (2025). Leading the green revolution: How transformational and green servant leadership drive employee sustainability in hospitality and tourism industry. *Indus Journal of Social Sciences*, 3(1), 420-445.
- Kalff, Y., & Simbeck, K. (2025). Explained, yet misunderstood: How AI literacy shapes HR managers' interpretation of user interfaces in recruiting recommender systems.
- Kamran, A., Syed, N. A., Ahmed, O., & Hussain, A. (2024, August). Does e-procurement adoption help textile industry of Pakistan for recruitment. In International Conference on Management Science and Engineering Management (pp. 1202-1228). Springer Nature Singapore.
- Kekez, I., Lauwaert, L., & Begičević Ređep, N. (2025). Is artificial intelligence (AI) research biased and conceptually vague? A systematic review of research on bias and discrimination in the context of using AI in human resource management. *Technology in Society*, 81, 102818. <https://doi.org/10.1016/j.techsoc.2025.102818>
- Köchling, A., & Wehner, M. C. (2020). Discriminated by an algorithm: A systematic review of discrimination and fairness by algorithmic decision-making in the context of HR recruitment and HR development. *Business Research*, 13, 795-848.
- Long, D., & Magerko, B. (2020). What is AI literacy? Competencies and design considerations. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems* (pp. 1-16). ACM.
- Moritz, J. M., & Schmidt, C. (2024). Trust in algorithmic management: The role of task type and prior discrimination experience. *Academy of Management Proceedings*, 2024(1). <https://doi.org/10.5465/AMPROC.2024.14136abstract>
- Mubashir, A., Arif, S., Kazmi, Q. A., Naqvi, S. T. A., & Ahmed, O. (2023). Stress management in the banking industry of Pakistan: Emotion regulation as a moderator in employee conflict behavior. *Journal of Education and Social Studies*, 4(2), 394-411.

- Muniza Syed, Ahmed, O., Azher, E., Salman, S., Siddiqui, H. M. A., & Javed, S. (2025). The impact of influencer marketing on consumer purchase intention: The mediating role of trust, content, consumer engagement, and popularity. *ASSA Journal*, 3(1), 147–166.
- Naoum, R. F., Szakadti, T., & Balogh, G. (2026). Artificial intelligence (AI) in human resource management (HRM): A systematic review of its dual impact on diversity, equity, and inclusion (DEI). *Management Review Quarterly*. <https://doi.org/10.1007/s11301-025-00580-y>
- Ng, D. T. K., Leung, J. K. L., Chu, S. K. W., & Qiao, M. S. (2021). Conceptualizing AI literacy: An exploratory review. *Computers and Education: Artificial Intelligence*, 2, 100041.
- Nguyen, T., & Connolly, R. (2025). Employee trust in artificial intelligence (AI)-assisted HRM: A performance evaluation context. *Academy of Management Proceedings*, 2025(1).
- Nyhan, R. C., & Marlowe, H. A. (1997). Development and psychometric properties of the organizational trust inventory. *Evaluation Review*, 21(5), 614–635.
- OSAMA, A., SADIA, J., EINAS, A., BUSHRA, J., & FARRUKH, Z. (2025). Leading the green revolution: How transformational and green servant leadership drive employee sustainability in hospitality and tourism industry. *INDUS JOURNAL OF SOCIAL SCIENCES* Учредители: Ali Institute of Research & Skills Development, 3(1), 420-445.
- Qin, S., Jia, N., Luo, X., Liao, C., & Huang, Z. (2023). Perceived fairness of human managers compared with artificial intelligence in employee performance evaluation. *Journal of Management Information Systems*, 40(4), 1039–1070.
- Robinson, S. L. (1996). Trust and breach of the psychological contract. *Administrative Science Quarterly*, 41(4), 574–599.
- Sadia, J., Iraj, F., Einas, A., Muniza, S., Osama, A., & Farrukh, Z. (2025). Bridging Learning and Technology: How Digital Platforms Impact Academic Performance. *Journal for Social Science Archives*, 3(1), 724-743.
- Siddiqui, H. M. A., Azher, E., Ali, O., Khan, M. F. U., Salman, S., & Ahmed, O. (2025). Synergizing Resilience and Digitalization: A New Paradigm in Supply Chain Performance. *The Critical Review of Social Sciences Studies*, 3(1), 3506-3526.
- Starke, C., Baleis, J., Keller, B., & Marcinkowski, F. (2021). Fairness perceptions of algorithmic decision-making: A systematic review of the empirical literature. *Big Data & Society*, 9(2).
- Strohmeier, S., Becker, M., & Scheer-Weller, E. (2026). Beyond aversion: Principles of appropriate algorithmic decision-making in human resource management. *Expert Systems with Applications*. <https://doi.org/10.1016/j.eswa.2025.128954>
- Wagan, S. M., & Sidra, S. (2025). Artificial intelligence in human resource management: A systematic review of adoption, impact, and challenges. *AYBU Business Journal*, 5(2).
- Wang, R., Harper, F. M., & Zhu, H. (2020). Factors influencing perceived fairness in algorithmic decision-making: Algorithm outcomes, development procedures, and individual differences. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*. Association for Computing Machinery.
- Zafar, F., Jabbar, B., Bano, A., Khan, M. F. U., Ahmed, O., & Salman, S. (2025). Nexus of univariate and multivariate models for forecasting interest rates in Pakistan using VAR–COV analysis. *Advance Journal of Econometrics and Finance*, 3(1), 68–84.
- Zhao, J., & Shi, J. (2026). How does AI literacy empower decision-making quality? The mediating role of employee-AI collaboration and the moderating role of AI trust. *Acta Psychologica*, 268, 107295.
- Zhou, J., Verma, S., Mittal, M., & Chen, F. (2021). Understanding relations between perception of fairness and trust in algorithmic decision making. *arXiv*. <https://doi.org/10.48550/arXiv.2109.14345>

