



JOURNAL OF SOCIAL SCIENCES DEVELOPMENT

www.jssd.org.pk



editor@jssd.org.pk

ALGORITHMIC TRANSPARENCY, PERCEIVED FAIRNESS & EMPLOYEE TRUST
IN AI-DRIVEN HR DECISIONS: A MEDIATION MODEL

Muhammad Farhan Khan¹ & Irfan Ullah Khan²

¹PhD Management Studies, Department of Public Administration, Gomal University, Dera Ismail Khan

²Associate Professor, Department of Public Administration, Gomal University, Dera Ismail Khan

KEYWORDS	ABSTRACT
Algorithmic Transparency; Perceived Fairness; Employee Trust; AI-Driven HR Decisions; Mediation Analysis	The algorithmic transparency in AI-driven HR decisions helps employees understand how decisions are made, reducing uncertainty and increasing perceptions of procedural fairness. Perceived fairness acts as a key mediator by translating the transparent AI processes into stronger employee trust in automated recruitment, evaluation, promotion, and reward decisions. This study examines whether algorithmic transparency builds employee trust in AI-driven HR decisions directly or indirectly through perceived fairness. Drawing on organizational justice theory and trust-in-automation theory, a cross-sectional survey of 322 mid-to-senior managers at the multinational corporations in Pakistan was analyzed using Hayes' PROCESS Macro (Model 4) with 5,000 bootstrap resamples. Transparency was positively associated with both trust ($\beta = 0.217, p < .001$) and perceived fairness ($\beta = 0.351, p < .001$), and fairness was positively associated with trust ($\beta = 0.295, p < .001$). The indirect effect of transparency on trust through fairness was significant ($\beta = 0.103, 95\% \text{ CI } (0.067, 0.142)$), with the direct effect remaining significant, signifying partial mediation. Results show transparency alone is insufficient to secure the employee trust; its impact depends on whether decisions are perceived as fair, with the implications for designing more trustworthy AI-driven HR systems.
ARTICLE HISTORY	
Date of Submission: 04-04-2026 Date of Acceptance: 12-05-2026 Date of Publication: 14-05-2026	
	 2026 Journal of Social Sciences Development
Corresponding Author	Muhammad Farhan Khan
Email:	farhanphd1@gmail.com
DOI	https://doi.org/10.53664/JSSD/05-02-2026-04-40-52

INTRODUCTION

Artificial Intelligence (AI) has grown from a passive analytics solution to a pro-active decision-making agent in Human Resource Management (HRM) in industries like recruitment, evaluation, promotion and resource allocation policy (Bankins, Formosa, Griep & Richards, 2022; Chen, 2026; Lima & Rahman, 2024). The employees perceive these automated HR decisions as social and moral

events and consider how fair and trustworthy they are, as they have a critical impact on their livelihoods and career paths. (Colquitt, Hill & Cremer, 2022; Newman, Fast & Harmon, 2020). Accordingly, the question of how to ensure employee trust in algorithmic decision-making (ADM) has emerged as key research topic in organizations (Starke, Baleis, Keller & Marcinkowski, 2021; Wang, 2025). It is often assumed in industry and initial literature that transparency is the single most important determinant for trust in algorithms: by exposing how the system works it is supposed to lead to legitimacy and trust (Lepri, Oliver, Letouzé, Pentland & Vinck, 2017; Grimmelikhuijsen, 2022). In this connection, the present study is the basis of the explainable-AI (XAI) efforts in HR technology (Fabeyo, 2025; Goyal, 2025). But empirical evidence shows a problematic path from transparency to trust.

The effect of transparency on informational justice and trust has been mixed; it can be beneficial in some cases (Ning, Lu, Li & Gupta, 2024; Ochmann, Michels, Maier & Laumer, 2024) but null or even counterproductive in others due to increased scrutiny or a lack of consideration for the human context, human effort and extenuating circumstances (Niessen & Neumann, 2026; Yu & Li, 2022). Mirbabaie, Langer, Rieskamp and Hofeditz (2025) call this 'transparency fallacy' since disclosure won't work if the system is seen as reductionist and/or values misaligned. This relationship reveals that transparency does not directly influence trust but rather it goes over intervening mechanism: the perceived fairness (Afroogh, Akbari, Malone, Kargar & Alambeigi, 2024; Shin & Park, 2019; Starke et al., 2021). Results of relative studies suggest that AI decisions generally lack procedural fairness because they are less transparent in terms of information and are made based on fairness instead of evaluation (Jialiang, Shanshi, Chuyan & Yifan, 2021). Likewise, in financial-inclusion and algorithmic recruitment domains, it has been concluded that transparency increases the sense of the fairness leading toward trust, satisfaction, and dependency (Yang & Lee, 2024; Zhou, Verma, Mittal & Chen, 2021).

Therefore, the proximal mechanism of the distal design characteristics, such as transparency, is that of emotional legitimacy (Wang, 2025; Yu & Li, 2022). Organizational justice theory views fairness as multi-faceted perception that includes distributive, procedural, interpersonal and informational justice for realizing the potential consequences (Cohen & Spector, 2001; Colquitt, Hill & Cremer, 2022). Thus, this framework provides an explanation for why formal computational fairness (e.g., statistical bias reduction) can be felt to be unfair by the employees when the qualitative context or human discretion is taken out, the phenomenon that has been coined 'algorithmic reductionism' (Köchling & Wehner, 2020; Newman, Fast & Harmon, 2020). In line with the theory of trust in automation, trust serves as the basis for relying on opaque systems (Lee & See, 2004; Shin, 2020). Transparency can only be a means for trust if explanations can be informative and build actual trust as well as not just the false sense of it (Angerschmid, Zhou, Theuermann, Chen & Holzinger, 2022; Grimmelikhuijsen, 2022). The Perceived fairness & trust are the foundation of the bridge between algorithmic design, organizational acceptance as well as behaviors (Kordzadeh & Ghasemaghahi, 2021; Wang, 2025).

The empirical studies focus mostly on the recruitment process, neglecting other processes such as performance and promotion to a great extent, where the fairness perceptions are also much stronger

(Lima & Rahman, 2024; Park, Ahn, Hosanagar & Lee, 2022). Second, many of the previous studies only focus on the bivariate relationships but not on testing a full mediation model in organizational context (Wang, 2025). The use of vignette experiments, western samples restrict generalizability at other cultural contexts where cultural norms & power distance affect transparency interpretation (Aysolmaz, Müller & Meacham, 2023; Yu & Li, 2022). Women's participation in the use of AI in the workplace is underscored by contextual boundary conditions, including human involvement in the task and its type (Bankins et al., 2023; Lee, 2018). This study attempts to fill gaps by proposing and testing a mediation model where algorithmic transparency has an indirect positive impact on employee trust via perceived fairness with implications of treating perceived fairness as the central psychological link to the algorithmic legitimacy (Angerschmid et al., 2022; Bankins et al., 2023; Rodgers et al., 2022).

LITERATURE REVIEW

Companies are increasingly incorporating AI systems into essential aspects of HR. These decisions have major implications for both an employee's personal and professional life and trust is the vital consideration (Bankins et al., 2022; Chen, 2026; Lima & Rahman, 2024). Literature indicates that transparency is not automatically equated with trust, but rather impact is dependent on employee fairness assessments (Afroogh et al., 2024; Ning et al., 2024; Ochmann et al., 2024; Shin & Park, 2019; Starke et al., 2021). This section reviews and integrates relevant theory and evidence on the algorithmic transparency, employee perceptions of fairness & trust in algorithms into our proposed mediation model.

Algorithmic Transparency

The transparency aspect of algorithms refers to degree to which logic, data & criteria of automated decision-making made transparent to users affected by these decisions (Grimmelikhuijsen, 2022; Jabagi, Croteau, Audebrand & Marsan, 2024; Yang & Lee, 2024). In HR it ranges from disclosure of a simple system to explanations provided on a granular level with regard to decision weights and variables (Niessen & Neumann, 2026; Ochmann et al., 2024). There is division amid transparency (information access), explainability (Fabeyo, 2025). Empirical evidence on benefits of transparency is inconclusive. The research indicates that process disclosure improves competence beliefs and that they increase the use of advice (Angerschmid et al., 2022); research on transparency in recruitment has not consistently proven to improve procedural or distributive justice. The transparency mainly impacts on level of distributive over informational fairness in platform work, which was mediated by prior tech attitudes (Mirbabaie et al., 2025). Besides, an easy disclosure can trigger mistrust and compel an investigation (Niessen & Neumann, 2026), and can lead to the 'transparency fallacy' if there are features in the system that seem reductionist or misaligned values (Newman et al., 2020; Schoeffer et al., 2021).

Perceived Fairness

Perceived Fairness refers to the subjective assessment of fairness of decision processes and outcomes, which are felt to be equitable, respectful and explanations were suitable (Colquitt, 2001; Colquitt et al., 2022; Ochmann et al., 2024). It is different from computational fairness (formal statistical bias reduction), since algorithms technically free from bias can still be perceived as unfair, if employees

think that there is a lack of human discretion or contextual nuances (Köchling & Wehner, 2020; Newman et al., 2020; Yu & Li, 2022). The justice model proposed by Colquitt (2001), consists of distributive, procedural, interpersonal and informational justice, is the predominant model in the literature on organizational fairness. Thus, does the algorithm rate lower than humans in terms of procedural and interactional fairness when applied to AI? This is frequently the case for complex, qualitative evaluative tasks (Bankins et al., 2022; Jialiang et al., 2021; Köchling & Wehner, 2020; Lee, 2018). Thus, it is employees' subjective perception of fairness that serves as the proximal engine driving links from system design to employee reactions (Bankins et al., 2023; Decker et al., 2024; Rodgers et al., 2022).

Employee Trust

Trust may be defined as the willingness to be vulnerable to the other party, based on an expectation of their competence, benevolence and integrity (Mayer et al., 1995). Trust, in the context of AI, is a measure of confidence in system's ability and reliability, influencing the extent to which users are willing to trust AI-generated results (Lee & See, 2004; Shin, 2020). System fairness is predictor of system trust (Afroogh et al., 2024; Shin, 2020; Starke et al., 2021; Zhou et al., 2021). Starke et al., (2021) report that downstream, trust and fairness influence organizational commitment, job pursuit intentions, organizational support as well as litigation risks (Irvan et al., 2026; Jabagi et al., 2024; Mulyatini et al., 2026).

Theoretical Framework

The organizational justice theory, and the theory on trust in automation, are integrated within our model. The organizational justice theory offers a multidimensional perspective that can be used to explain the success or failure of technical transparency in inducing a sense of legitimacy (Ochmann et al., 2024; Colquitt, 2001; Yang & Lee, 2024). When the system is opaque, there is a distinction between its ability and its predictability which affects the level of dependence of the user on the system, based on trust-in-automation theory (Lee & See, 2004; Mayer, Davis, & Schoorman, 1995). Together, these theories prescribe an algorithmic design (transparency) towards justice judgments (perceived fairness) and, in turn, towards the employees' trust (Wang, 2025; Yang, Lu & Cooke, 2025; Yu & Li, 2022).

Relationship between Algorithmic Transparency & Employee Trust

Transparency theory argues that transparency will lead to decrease in opacity, system competence beliefs and dependence (Grimmelikhuijsen, 2022; Lee & See, 2004; Ning et al., 2024; Shin, 2020). Empirical studies conducted recently support positive links amid trust and transparent practices of AI (Ningsih et al., 2026; Warsono et al., 2025). But inconclusive results in the cases of simple hiring or gigs indicate that direct dependency cannot be taken as a given (Mirbabaie et al., 2025; Niessen & Neumann, 2026).

H1: Algorithmic transparency is positively related to employee trust in the AI-driven HR decisions.

Relationship between Algorithmic Transparency & Perceived Fairness

Transparency serves as informational cue that employees can use to determine the accountability and respectfulness of decision-making process (Grimmelikhuijsen, 2022; Jabagi et al., 2024; Yang

& Lee, 2024). These revelations help to improve informational justice in recruitment and more widespread algorithmic settings (Aysolmaz et al., 2023; Ochmann et al., 2024; Yang & Lee, 2024). Still, transparency is not necessarily equated with increases in procedural justice and distributive justice, as it is the quality of the information made transparent that determines the justice outcome (Niessen & Neumann, 2026).

H2: AT is positively related to employees' perceived fairness of AI-driven HR decisions.

Relationship between Perceived Fairness and Employee Trust

Fairness perception is a fundamental precursor to the trust in the organization and system (Cohen-Charash & Spector, 2001; Starke, Baleis, Keller & Marcinkowski, 2021). The experimental and review results show that positive fairness cues always increase trust while negative ones (perceived discrimination) decrease trust in the automated systems (Afroogh et al., 2024; Zhou, Verma, Mittal & Chen, 2021).

H3: The perceived fairness is positively related to employee trust in the AI-driven HR decisions.

Perceived Fairness as a Mediator

In other words, synthesizing these connections, perceived fairness serves as the major mechanism to explain the relationship between transparency and trust (Jialiang et al., 2021; Yang & Lee, 2024). Mediation helps explain variation in direct transparency-trust relationships; transparency helps to foster the trust only if it is perceived as fair and legitimatizing environments (Niessen & Neumann, 2026; Yu & Li, 2022; Mirbabaie, Langer, Rieskamp & Hofeditz, 2025). This mediated path is influenced by limit conditions like type of task, mechanical vs human-skill; human involvement, hybrid human-AI decision making & individual technology attitude (Araujo et al., 2020; Bankins et al., 2023; Lee, 2017, 2018; Lee, Jain, Cha, Ojha & Kusbit, 2019; Park, Ahn, Hosanagar & Lee, 2022; Tehreem, 2025).

H4: The PF mediates relationship between AT and employee trust in the AI-driven HR decisions.

Research Gap

Despite increased attention, there are three major gaps: (1) overly focused upon recruitment to the exclusion of performance, promotion and compensation (Lima & Rahman, 2024); (2) fragmented testing of separate bivariate links in field settings (Ochmann et al., 2024; Starke et al., 2021; Wang, 2025; Yang & Lee, 2024); and (3) reliance on the Western experimental vignettes limiting cross-cultural generalizability (Aysolmaz et al., 2023; Starke et al., 2021; Yu & Li, 2022). In this context, this study aims to fill this gap by examining "combined mediation" model for a variety of AI-based human resource decisions, building upon the justice-based theoretical models of organizational AI governance (Angerschmid et al., 2022; Bankins et al., 2023; Chen, 2026; Fabeyo, 2025; Köchling & Wehner, 2020).

RESEARCH METHODOLOGY

The research approach used in this study was quantitative, cross-sectional survey, using a survey method to carry out the study in an empirical manner to test the proposed theoretical relationships between algorithmic transparency, perceived fairness, and employee trust in AI decision-making in human resource management (HRM). A cross-sectional design would be appropriate to obtain a

snapshot of the current organizational processes and to measure the perceptual models of current organizational members (Spector, 2019). This research design enabled the simultaneous evaluation of the main independent variable (Algorithmic Transparency), the mediator (Perceived Fairness), and the key outcome.

Target Population & Sampling Technique

The target population was the mid-to-senior level managers employed in leading multinational corporations (MNCs) with AI-driven HRM and algorithmic decision-making practices in Pakistan. Managers have been selected due to nature of work where they are exposed to automated solutions that are used for talent acquisition, performance evaluations and promotion shortlisting, and task allocation. A non-probability (purposive) sampling was used. To ensure that the respondents had relevant and firsthand experience with AI-supported HR tools and were working in agile, tech-centric work environments, sampling was focused on this group. Five big MNCs in various industries of Pakistan: Unilever Pakistan, Nestlé Pakistan, Coca-Cola Pakistan, Abbott Laboratories Pakistan, and PepsiCo Pakistan. The selection of particular organizations was driven by three considerations: they are foreign-based, they operating in stable manner, and they have large pools of managerial talent that have gained first-hand experience with corporate digital transformation. Compared to domestic business, MNCs have greater competitive environments globally and are the first to adapt unified algorithmic HR.

Sample Size & Data Collection Procedure

Structured self-administered questionnaire was used for data collection. Five MNCs distributed a total of 365 questionnaires among human resource and managerial department of each of the five MNCs. The participation was voluntary and the respondents were guaranteed confidentiality and anonymity of the data. Thus, after collecting the data, instruments were checked for missing data, unengaged response pattern, and outliers. A total of 322 surveys were received and deemed usable for final statistical analysis reflecting an effective response rate of 88.2% of the distributed surveys. This number of observations meets the overall guidelines of at least 10 observations per estimated parameter that are recommended for general structural equation modeling and regression analysis (Hair et al., 2019).

Operationalization & Measurement Instruments

All constructs were measured with previous scales that were adapted to the context of algorithmic HR decision-making, and all had been previously validated. The survey items were rated using a 5-point standard Likert scale from strongly disagree to strongly agree. Algorithmic Transparency: A 3-item scale adapted from Yagmurdur et al. (2026). The items reflect the extent of employees' perceptions of algorithms' logic, criteria, decision making as clear, understandable, and accessible. Perceived Fairness: It was assessed with a 4-item scale derived from Yagmurdur et al. (2026). The fairness perceptions of processes and results produced by algorithmic decisions (The decision made by the algorithm was fair"). Employee Trust in AI-Driven HR Decisions (Dependent Variable): This study adapted an 11-item scale from Do et al. (2025) for this variable. The items reflect employees' confidence, dependency and positive attitudes to use of AI in the organization ("I have confidence

in use of AI technology" and "I believe AI technology will consistently operate providing adequate and efficient results").

DATA ANALYSIS

The SPSS Statistics software and Hayes PROCESS Macro (Version 4.2) were used for the statistical analyses. Analytical procedure was carried out in accordance with the four stages analysis process towards the conclusion.

Table 1 Descriptive Statistics & Internal Consistency Reliability (N = 322)

Construct	No. of Items	Mean	SD	CA (α)
Algorithmic Transparency (AT)	3	3.557	0.717	0.729
Perceived Fairness (PF)	4	3.436	0.665	0.811
Employee Trust in AI (TAI)	11	3.651	0.550	0.857

The mean and standard deviation for the algorithmic transparency were 3.557 (0.717) and 0.729 respectively as reported in Table 3.1. The mean of perceived fairness was 3.436 (SD = 0.665) and the reliability was good ($\alpha = 0.811$). similarly, the employee trust in AI yielded a mean of 3.651 (SD = 0.550) and an alpha of 0.857. in this linking, all the constructs have very good reliability with the ranges of 0.70–0.87.

Construct Validity and Factor Structure

The construct validity and factor structure confirm that the measurement items precisely represent their intended constructs and load appropriately upon their respective factors. In order to check structural validity, an Exploratory Factor Analysis (EFA) was done using the Principal Component Analysis (PCA) with Varimax orthogonal rotation for all 18 items. Three factors were extracted and the analysis revealed three-factor solution with eigenvalues greater than 1.0, which explained 61.4% of cumulative variance. The results of Harman's single factor test came out at 34.2% of the total variance, which is less than the 50% threshold and henceforth there is no Common Method Variance (CMV) contamination of study results. The primary factor loadings were all above the recommended value of 0.5 and cross-loadings were low (< 0.30), indicating high convergent and divergent validity.

Table 2 Pearson Correlation Matrix

Constructs	[1]	[1]	[1]
Algorithmic Transparency [1]	1.000	—	—
Perceived Fairness [1]	0.378***	1.000	—
Employee Trust in AI [1]	0.283***	0.413***	1.000

Note: *** $p < 0.001$ (two-tailed).

The correlation matrix shows that there are positive linear relationships. Algorithmic Transparency has significant positive correlation with perceived fairness ($r = 0.378, p < 0.001$) and employee trust in AI ($r = 0.283, p < 0.001$). There is high correlation amid perceived fairness & employee trust in AI ($r = 0.413, p < 0.001$).

Table 3 Direct, Indirect & Total Effects (Hayes PROCESS Model 4)

Hypothesis / Path	(β)	SE	t	p	LLCI	ULCI	Decision
H1: AT $\rightarrow$ Trust (Total)	0.217	0.041	5.276	< 0.001	0.136	0.298	Supported
H2: AT $\rightarrow$ PF (Path a)	0.351	0.048	7.308	< 0.001	0.256	0.445	Supported
H3: PF $\rightarrow$ Trust (Path b)	0.295	0.045	6.553	< 0.001	0.207	0.384	Supported
AT $\rightarrow$ Trust (Direct c')	0.113	0.042	2.716	0.007	0.031	0.196	Supported
H4: Indirect (AT $\rightarrow$ PF $\rightarrow$ Trust)	0.103	0.019	—	—	0.067	0.142	Supported

Interpretation of Findings

The results show significant overall positive impact of algorithmic transparency on employee trust in AI decisions (H1) ($\beta = 0.217$, $t = 5.276$ and $p < 0.001$). So, hypothesis H1 is accepted. Model H2 – Path a show that algorithmic transparency has a high impact on perceived fairness for managers ($\beta = 0.351$, $t = 7.308$, $p < 0.001$) accounting for 14.3% of the variance in fairness perceptions ($R^2 = 0.143$). There is sufficient evidence to support hypothesis H2. 3. Perceived Fairness to Employee Trust (H3 – Path b): There is a strong direct effect of perceived fairness to employee trust in AI decision making ($\beta = 0.295$, $t = 6.553$, $p < 0.001$). Hypothesis H3 is completely accepted. 4. Mediation Effect of the Perceived Fairness (H4): The indirect bootstrap effect of the algorithmic transparency on trust (via perceived fairness) was $\beta = 0.103$ ($SE = 0.019$). The results show that bootstrap 95% CI for the indirect effect does not contain zero (CI values: [0.067, 0.142]) and this indicates a statistically significant indirect effect. The direct effect ($c' = 0.113$, $p = 0.007$) is still present after the addition of fairness to the model, thereby indicating partial mediation of Perceived Fairness. There is complete support for hypothesis H4.

DISCUSSION

The findings from this study provide compelling evidence for all four hypotheses and when viewed as a whole, bring to light relationship that has been addressed in an inconsistent manner in previous literature. In line with H1, algorithmic transparency had a positive relationship with employee trust in AI-powered HR decisions. This aligns with recent research indicating providing transparency regarding the algorithm logic can improve trust in automated decision-making. The size of this total effect ($\beta = 0.217$) was small, however, in comparison to the effect of perceived fairness on trust ($\beta = 0.295$, H3), and the mediation analysis (H4) indicated that a significant amount of the effect of transparency happens indirectly via the mediator of perceived fairness, as opposed to directly via transparency. This practice provides empirical support for the 'transparency fallacy' argument presented in recent research: mere transparency on functioning of algorithms does not lead to trust, if employees do not perceive the process and products as the equitable, respectful and sensitive to human context. The findings reveal that while transparency has some residual direct value ($c' = 0.113$) after controlling for fairness, there is an indirect effect: fairness perceptions are a closer proxy of the direct impact of transparency, such as in the form of organizational accountability, as opposed to disclosure itself.

The results of the study have both theoretical and practical implications. In theory, the study builds on organizational justice theory and trust-in-automation theory to the field of the algorithmic HR

management and shows that its effect on trust can only be explained through a more proximal mechanism, namely fairness judgements: transparency acts as a distal design factor that is unlikely to have a direct influence on trust and is therefore likely to work through a more proximal one, in this case, fairness judgements in diverse circumstances. This comprehensive framing also aligns with technology-acceptance perspectives, indicating that if AI-powered HR tools are not only useful and easy to use, but also socially and cognitively appraised as fair, employees will likely be more inclined to accept them. Thus, as the sample comprised mid-to-senior managers in multinational companies in the single national context and data were obtained via a cross-sectional self-report survey, the findings should be interpreted with caution, and further studies with longitudinal or multi country designs could help to confirm the durability and generalization of the mediated path found in this study.

CONCLUSION

This research argues that the algorithmic transparency is an essential building block to developing employee trust in AI decision mechanisms within HR. In addition, there is a necessary explanatory gap between perceived fairness and the rest. With transparent and proven fairness at the core of automated HR tools, organizations can build trust in the workforce, drive digital HR transformation, and improve operational effectiveness in agile organizational models. From a practical perspective, the findings indicate that the transparency in HR processes should be complemented by visible measures to ensure the fairness, such as options for the human intervention, contextually adjusted explanations, and opportunities to challenge automated decisions, especially in more subjective HR-related activities such as performance ratings as well as promotion. Consequently, the greater algorithmic transparency enhances employee trust both directly and indirectly through improved perceived fairness.

Theoretical Implications

This research contributes to the theoretical understanding in three main dimensions based on the empirical investigation, which is grounded in theories of Organizational Justice Theory (OJT) and the Technology Acceptance Model (TAM): The TAM typically argues that utilitarian perceptions of the one's system use (Perceived Usefulness (PU) and Perceived Ease of Use (PEOU)) are the most important factors to consider, but it fails to account for socio-cognitive and fairness considerations employees have when engaging with the opaque and autonomous AI decision-making tools. This study demonstrates that the OJT and TAM are not mutually exclusive, as combining OJT and TAM moves technology acceptance beyond just usability factors, shows that both cognitive appraisals of TAM (informational and procedural) are essential to consider in the evaluation of the employees' trust in AI systems.

REFERENCES

Afroogh, S., Akbari, A., Malone, E., Kargar, M., & Alambeigi, H. (2024). The Trust in AI: Progress, challenges, and future directions. *Humanities and Social Sciences Communications*, 11. <https://doi.org/10.1057/s41599-024-04044-8>.

- Angerschmid, A., Zhou, J., Theuermann, K., Chen, F., & Holzinger, A. (2022). Fairness and explanation in AI-informed decision making. *Machine Learning and Knowledge Extraction*, 4, 556–579. <https://doi.org/10.3390/make4020026>.
- Araujo, T. B., Helberger, N., Kruijkemeier, S., & de Vreese, C. D. (2020). In AI we trust? Perceptions about automated decision-making by artificial intelligence. *AI & Society*, 35, 611–623. <https://doi.org/10.1007/s00146-019-00931>.
- Aysolmaz, B., Müller, R., & Meacham, D. (2023). The public perceptions of algorithmic decision-making systems: Results from a large-scale survey. *Telematics and Informatics*, 79, 101954. <https://doi.org/10.1016/j.tele.2023.101954>.
- Bankins, S., Formosa, P., Griep, Y., & Richards, D. (2022). AI decision making with dignity? Contrasting workers' justice perceptions of human and AI decision making in a human resource management context. *Information Systems Frontiers*, 24, 857–875. <https://doi.org/10.1007/s10796-021-10223-8>.
- Bankins, S., Ocampo, A. C., Marrone, M., Restubog, S. L. D., & Woo, S. E. (2023). A multilevel review of artificial intelligence in organizations: Implications for organizational behavior research and practice. *Journal of Organizational Behavior*. <https://doi.org/10.1002/job.2735>.
- Chen, Z. (2026). Algorithmic human resource management as a mode of algorithmic governance: Transparency, fairness, and human agency in the digital workplace. *Humanities and Social Sciences Communications*. <https://doi.org/10.1057/s41599-026-06989-4>.
- Cohen-Charash, Y., & Spector, P. E. (2001). The role of justice in organizations: A meta-analysis. *Organizational Behavior and Human Decision Processes*, 86(2), 278–321. <https://doi.org/10.1006/obhd.2001.2958>.
- Colquitt, J. A. (2001). On the dimensionality of organizational justice: A construct validation of a measure. *Journal of Applied Psychology*, 86(3), 386–400. <https://doi.org/10.1037/0021-9010.86.3.386>.
- Colquitt, J., Hill, E. T., & Cremer, D. (2022). Forever focused on fairness: 75 years of organizational justice in Personnel Psychology. *Personnel Psychology*. <https://doi.org/10.1111/peps.12556>.
- Decker, M., Wegner, L., & Leicht-Scholten, C. (2024). Procedural fairness in algorithmic decision-making: The role of public engagement. *Ethics and Information Technology*, 27. <https://doi.org/10.1007/s10676-024-09811-4>.
- Do, H., Chu, L. X., & Shipton, H. (2025). How and when AI-driven HRM promotes employee resilience and adaptive performance: A self-determination theory. *Journal of Business Research*, 192, 115279. <https://doi.org/10.1016/j.jbusres.2025.115279>.
- Fabeyo, S. (2025). Explainable AI in employment decision-making: A systematic review of transparency methods in hiring algorithms. *Issues in Information Systems*. https://doi.org/10.48009/3_iis_2025_2025_110.
- Grimmelikhuijsen, S. (2022). Explaining why the computer says no: Algorithmic transparency affects the perceived trustworthiness of automated decision-making. *Public Administration Review*. <https://doi.org/10.1111/puar.13483>.
- Hair, J. F., Black, W. C., Babin, B. J., & Anderson, R. E. (2019). *Multivariate data analysis* (8th ed.). Cengage Learning.

- 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*. <https://doi.org/10.59431/ijer.v6i1.729>.
- Jabagi, N., Croteau, A., Audebrand, L. K., & Marsan, J. (2024). Do algorithms play fair? Analyzing the perceived fairness of HR-decisions made by algorithms and their impacts on gig-workers. *The International Journal of Human Resource Management*, 36, 235–274. <https://doi.org/10.1080/09585192.2024.2441448>.
- Jialiang, P., Shanshi, L., Chuyan, Z., & Yifan, C. (2021). Can AI algorithmic decision-making improve employees' perception of procedural fairness? *Foreign Economics & Management*, 41–55. <https://doi.org/10.16538/j.cnki.fem.20210818.103>.
- 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*. <https://doi.org/10.1007/s40685-020-00134-w>.
- Kordzadeh, N., & Ghasemaghaei, M. (2021). Algorithmic bias: Review, synthesis, and future research directions. *European Journal of Information Systems*, 31, 388–409. <https://doi.org/10.1080/0960085x.2021.1927212>.
- Lee, J. D., & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46(1), 50–80. <https://doi.org/10.1518/hfes.46.1.50.30392>.
- Lee, M. K. (2017). Algorithmic mediation in group decisions: Fairness perceptions of algorithmically mediated vs. discussion-based social division. Proceedings of the 2017 ACM Conference on Computer Supported Cooperative Work and Social Computing. <https://doi.org/10.1145/2998181.2998230>.
- Lee, M. K. (2018). Understanding perception of algorithmic decisions: Fairness, trust, and emotion in response to algorithmic management. *Big Data & Society*, 5(1). <https://doi.org/10.1177/2053951718756684>.
- Lee, M. K., Jain, A., Cha, H. J., Ojha, S., & Kusbit, D. (2019). Procedural justice in algorithmic fairness. Proceedings of the ACM on Human-Computer Interaction, 3(CSCW). <https://doi.org/10.1145/3359284>.
- Lepri, B., Oliver, N., Letouzé, E., Pentland, A., & Vinck, P. (2017). Fair, transparent, and accountable algorithmic decision-making processes. *Philosophy & Technology*, 31, 611–627. <https://doi.org/10.1007/s13347-017-0279-x>.
- Lima, S. A., & Rahman, M. M. (2024). Algorithmic fairness in HRM balancing AI-driven decision making with inclusive workforce practices. *Journal of Information Systems Engineering and Management*. <https://doi.org/10.52783/jisem.v9i4s.13251>.
- Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *Academy of Management Review*, 20(3), 709–734. <https://doi.org/10.2307/258792>.
- Mirbabaie, M., Langer, M., Rieskamp, J., & Hofeditz, L. (2025). Transparency fallacy: Perceived fairness in algorithmic management. *Business & Information Systems Engineering*, 68, 829–850. <https://doi.org/10.1007/s12599-025-00963-1>.
- Mulyatini, N., Iskandar, Y., Sosiady, M., Maria, E., & Cahyati, P. (2026). AI-driven human resource management and its impact on employee performance: The mediating role of employee trust

- in artificial intelligence. *Dinasti International Journal of Education Management and Social Science*. <https://doi.org/10.38035/dijemss.v7i4.6436>.
- Newman, D. T., Fast, N. J., & Harmon, D. J. (2020). When eliminating bias isn't fair: Algorithmic reductionism and procedural justice in human resource decisions. *Organizational Behavior and Human Decision Processes*, 160, 149–167. <https://doi.org/10.1016/j.obhdp.2020.03.008>.
- Niessen, A., & Neumann, M. (2026). Transparency and trust in simple algorithmic hiring procedures. *International Journal of Selection and Assessment*. <https://doi.org/10.1111/ijsa.70045>.
- Ning, X., Lu, Y., Li, W., & Gupta, S. (2024). How transparency affects algorithmic advice utilization: The mediating roles of trusting beliefs. *Decision Support Systems*, 183, 114273. <https://doi.org/10.1016/j.dss.2024.114273>.
- Ningsih, S., Syahputra, D., Husainah, N., Wibowo, A., & Hakim, C. (2026). The influence of AI ethics in HR and algorithm transparency on employee trust in the manufacturing industry. *West Science Social and Humanities Studies*. <https://doi.org/10.58812/wsshs.v4i03.2702>.
- Ochmann, J., Michels, L., Maier, C., & Laumer, S. (2024). Perceived algorithmic fairness: An empirical study of transparency anthropomorphism in algorithmic recruiting. *Information Systems Journal*, 34, 384–414. <https://doi.org/10.1111/isj.12482>.
- Park, H., Ahn, D., Hosanagar, K., & Lee, J. (2022). Designing fair AI in human resource management: Understanding tensions surrounding algorithmic evaluation and envisioning stakeholder-centered solutions. *Proceedings of 2022 CHI Conference on Human Factors in Computing Systems*. <https://doi.org/10.1145/3491102.3517672>.
- Podsakoff, P. M., MacKenzie, S. B., & Podsakoff, N. P. (2012). Sources of method bias in social science research and recommendations on how to control it. *Annual Review of Psychology*, 63(1), 539–569. <https://doi.org/10.1146/annurev-psych-120710-100452>.
- Rodgers, W., Murray, J., Stefanidis, A., Degbey, W. Y., & Tarba, S. (2022). An artificial intelligence algorithmic approach to ethical decision-making in human resource management processes. *Human Resource Management Review*. <https://doi.org/10.1016/j.hrmr.2022.100925>.
- Schoeffer, J., Machowski, Y., & Kuehl, N. (2021). A study on fairness and trust perceptions in automated decision making. ARXIV. <https://doi.org/10.48550/arxiv.2103.04757>.
- Shin, D. (2020). User perceptions of algorithmic decisions in the personalized AI system: Perceptual evaluation of fairness, accountability, transparency, and explainability. *Journal of Broadcasting and Electronic Media*, 64(4), 541–565. <https://doi.org/10.1080/08838151.2020.1843357>.
- Shin, D., & Park, Y. (2019). Role of fairness, accountability, and transparency in algorithmic affordance. *Computers in Human Behavior*, 98, 277–284. <https://doi.org/10.1016/j.chb.2019.04.019>.
- Spector, P. E. (2019). Do not cross me: Optimizing the use of cross-sectional designs. *Journal of Business and Psychology*, 34(2), 125–137. <https://doi.org/10.1007/s10869-018-9570-x>.
- 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. <https://doi.org/10.1177/20539517221115189>.

- Tehreem, S. (2025). Explainable AI in public policy: Quantifying trust and distrust in algorithmic decision-making across marginalized communities. *AI and Machine Learning Advances*. <https://doi.org/10.64220/amla.v1i1.004>.
- Warsono, H. Y., Widada, D., & Sumarliani, S. (2025). Managerial consequences of the algorithmic determination of incentive decisions on procedural justice perceptions in production line. *Asian Journal of Management Analytics*. <https://doi.org/10.55927/ajma.v4i3.14997>.
- Yagmurdur, G., Meng, Y., & Gayaker, S. (2026). The role of algorithmic anthropomorphism, transparency, and fairness in shaping consumer purchase intentions in e-commerce: Evidence from Türkiye. *Journal of Theoretical and Applied Electronic Commerce Research*, 21(5), 159–178.
- Yang, M. M., Lu, Y., & Cooke, F. L. (2025). Demystifying AI for the workforce: The role of explainable AI in worker acceptance and management relations. *Journal of Management Studies*. <https://doi.org/10.1111/joms.70039>.
- Yang, Q., & Lee, Y.-C. (2024). Ethical AI in financial inclusion: The role of algorithmic fairness on user satisfaction and recommendation. *Big Data and Cognitive Computing*, 8, 105. <https://doi.org/10.3390/bdcc8090105>.
- Yu, L., & Li, Y. (2022). Artificial intelligence decision-making transparency and employees' trust: The parallel multiple mediating effect of effectiveness and discomfort. *Behavioral Sciences*, 12. <https://doi.org/10.3390/bs12050127>.
- Zhou, J., Verma, S., Mittal, M., & Chen, F. (2021). Understanding relations between perception of fairness and trust in algorithmic decision making. 2021 8th International Conference on Behavioral and Social Computing (BESC), 1–5. <https://doi.org/10.1109/besc53957.2021.9635182>.