



International Journal of Business and Management Sciences

2026; 2(4): 87-90 | ISSN(O): 2693-3500

www.sciro.org/ijbms

DOI: [10.5281/zenodo.21972882](https://doi.org/10.5281/zenodo.21972882)

Job Relevance, Trialability and Accountability in Applicants' Adoption of AI Recruitment

Deepak Kumar N^{1,*}, Nayana R², Meghana MS³.

¹ Assistant Professor, School of Management, CMR University, Bengaluru.

^{2,3} PG Scholar, Master Business Administration, School of Management, CMR University, Bengaluru.

Email address*: deepak.n@cmr.edu.in

Received: 12.06.2026 | **Revised**: 23.07.2026 | **Accepted**: 04.08.2026

Abstract: This study examines why applicants accept artificial-intelligence-enabled recruitment systems and whether responsible-AI perceptions explain adoption beyond practical utility.

Cross-sectional questionnaire data were obtained from 181 applicants, including students, fresh graduates, employees, job seekers and recruiters. Nine predictors—fairness, accountability, transparency, complexity, anxiety, trust, insecurity, trialability and job relevance—were entered into a multiple-regression model. Scale reliability and model arithmetic were audited before interpretation.

The model explained 69.8% of adoption variance (adjusted $R^2 = .682$; $F(9, 171) = 43.916$, $p < .001$). Job relevance was the dominant predictor ($\beta = .471$, $p < .001$), followed by trialability ($\beta = .248$, $p = .003$), complexity ($\beta = .204$, $p = .002$), and accountability ($\beta = .183$, $p = .002$). Fairness, transparency, anxiety, trust, and insecurity were non-significant. Because item wording and scale direction are unavailable, the positive complexity coefficient is treated as an unresolved measurement issue rather than evidence that difficulty increases adoption.

The paper identifies utility–governance decoupling in algorithmic hiring: adoption may proceed through relevance and experimentation even when fairness and trust remain statistically dormant. Responsible recruitment therefore requires governance safeguards that are not contingent on applicants rewarding them through adoption.

Keywords: AI recruitment, applicant reactions, responsible AI, job relevance, trialability, accountability, fairness, adoption.

JEL classification: J24, M51, O33

1. Introduction

Artificial intelligence now mediates recruitment activities that were once handled almost entirely by people: advertising, sourcing, résumé screening, assessment, interview scheduling and candidate communication. These systems promise scale and consistency, yet they also influence access to employment. A rejected recommendation can affect livelihood before an applicant knows what data were used or how the decision was reached.

This dual character makes AI recruitment different from ordinary workplace software. Usefulness and convenience matter, but legitimacy also depends on fairness, transparency, privacy, contestability and human oversight. The European Union treats AI used for recruitment and selection as a high-risk application under the AI Act (European Union, 2024). In India, the Digital Personal Data Protection Act requires lawful and purpose-

bound processing of digital personal data (Government of India, 2023). The NIST AI Risk Management Framework similarly emphasises governing, mapping, measuring and managing AI risk (NIST, 2023).

Applicant-reaction research shows that automation can improve efficiency while reducing perceived procedural justice when human contact, explanation or opportunity to perform is constrained (Langer et al., 2019; van Esch et al., 2021). More recent recruitment studies emphasise job relevance, explainability, augmentation and the possibility of trial (Horodyski, 2023; Lee & Cha, 2023). Yet adoption models often assume that favourable ethical perceptions naturally produce willingness to use.

The present study tests that assumption with 181 respondents. Its empirical contribution is a paradox: job relevance, trialability and accountability predict adoption, but fairness, transparency and trust do not once the predictors are consid-

ered jointly. The paper argues that ethical safeguards may be normatively indispensable even when they do not function as marginal adoption incentives.

2. Theoritical Framework and Hypothesis

2.1. Practical utility: job relevance and trialability

Technology-adoption theory proposes that users accept systems that improve task performance and are feasible to use (Davis, 1989; Venkatesh et al., 2003). Job relevance captures the fit between AI recruitment and the applicant's search task. Trialability reduces uncertainty by allowing candidates to encounter, test or rehearse a system before consequential use (Rogers, 2003).

2.2. Responsible-AI governance

Fairness concerns equitable treatment and avoidance of discriminatory outcomes. Accountability identifies who owns a decision, audits the system and provides remedy. Transparency concerns meaningful information about data, logic and process. These attributes are related but not interchangeable: an explanation can reveal an unfair rule, while accountability can create recourse even when the model remains technically complex (Rigotti & Fosch-Villaronga, 2024).

2.3. Psychological response: trust, anxiety and insecurity

Trust represents willingness to rely on the system under uncertainty. Anxiety captures apprehension about interacting with AI, and insecurity reflects perceived vulnerability or loss of control. Applicant reactions depend on the stakes, type of assessment and availability of human review; broad attitudes may therefore have less predictive value than system-specific experience.

2.4. Complexity

Complexity usually reduces innovation adoption. However, the source questionnaire does not report item wording or whether higher scores denote difficulty, sophistication or recoverability. A directional hypothesis is retained from the original model, but any positive coefficient must be interpreted cautiously until scale orientation is verified.

2.5. Hypotheses

- H1: Perceived fairness is positively associated with AI recruitment adoption.
- H2: Perceived accountability is positively associated with AI recruitment adoption.
- H3: Perceived transparency is positively associated with AI recruitment adoption.

- H4: Perceived complexity is negatively associated with AI recruitment adoption.
- H5: AI-related anxiety is negatively associated with AI recruitment adoption.
- H6: Trust is positively associated with AI recruitment adoption.
- H7: Insecurity is negatively associated with AI recruitment adoption.
- H8: Trialability is positively associated with AI recruitment adoption.
- H9: Job relevance is positively associated with AI recruitment adoption.

3. Method

3.1. Design and participants

A descriptive, cross-sectional design was used. Convenience sampling produced 181 responses from a heterogeneous pool described as students, fresh graduates, employees, job seekers and recruiters across IT, manufacturing, service, food-and-beverage and other sectors. The source paper does not provide a complete demographic frequency table, recruitment locations or response rate; population representativeness cannot be assessed.

3.2. Measures

The instrument contained 38 substantive items: fairness (3), accountability (3), transparency (3), complexity (4), anxiety (5), trust (4), insecurity (3), trialability (3), job relevance (4) and AI recruitment adoption (3). All constructs used multi-item Likert-type responses. Item wording, response anchors and scale sources must be recovered before submission.

Table 1: Internal consistency of study scales

Construct	Items	Cronbach's α	Assessment
Fairness	3	.853	Good
Accountability	3	.799	Acceptable
Transparency	3	.825	Good
Complexity	4	.893	Good
Anxiety	5	.823	Good
Trust	4	.886	Good
Insecurity	3	.805	Good
Trialability	3	.891	Good
Job relevance	4	.917	Excellent
AI recruitment adoption	3	.879	Good

3.3. Analysis and integrity checks

SPSS 20 was used for reliability, group comparisons, and multiple regression. Regression arithmetic is internally consistent: $84.329/120.814 = .698$, and $9.370/.213 \approx 43.99$, closely matching the reported F . Assumptions concerning linearity, normality, homoscedasticity, independence, and multicollinearity were not reported. ANOVA results are treated as exploratory because degrees of freedom, group sizes, omnibus statistics, post-hoc tests, and effect sizes are missing.

4. Results

4.1. Model explanatory power

The nine-predictor model was statistically significant, $F(9,171) = 43.916$, $p < .001$. It explained 69.8% of observed adoption variance (adjusted $R^2 = .682$; standard error = .462). This is substantial explanatory power, but the common-source design and overlapping constructs may inflate it.

4.2. Predictor effects

Job relevance is the dominant predictor. Trialability, complexity and accountability also make statistically significant contributions. Fairness, transparency, anxiety, trust and insecurity do not add unique explanatory power. Non-significance does not demonstrate that these factors are unimportant; correlated predictors, broad measures or indirect effects may absorb their variance.

Table 2: Regression results and corrected hypothesis decisions

Hypothesis / predictor	B	β	t	p	Decision
H1 Fairness	-.095	-.098	-1.485	.139	Not supported
H2 Accountability	.181	.183	3.170	.002	Supported
H3 Transparency	-.075	-.074	-1.332	.185	Not supported
H4 Complexity (expected -)	.203	.204	3.084	.002	Opposite direction
H5 Anxiety	.011	.012	.260	.795	Not supported
H6 Trust	.014	.014	.183	.855	Not supported
H7 Insecurity	-.005	-.005	-.068	.946	Not supported
H8 Trialability	.232	.248	3.069	.003	Supported
H9 Job relevance	.476	.471	6.087	< .001	Supported

4.3. Exploratory group differences

Knowledge about AI was associated with higher reported fairness, accountability, transparency, trust, trialability and job relevance. Respondents with direct AI-recruitment experience often reported lower means than those answering “no” or “maybe,” suggesting a possible experience–expectation gap. These patterns require cautious follow-up because several source rows are contradictory—for example, identical displayed adoption means are paired with $p = .016$ —and category sizes are unknown.

5.3. The complexity anomaly

The positive complexity coefficient conflicts with standard diffusion theory. Three explanations are plausible: items may be reverse-coded; the construct may capture sophistication or recoverability rather than difficulty; or multicollinearity may reverse the partial coefficient. The result should not be theorised further until the questionnaire and zero-order correlations are examined.

5. DISCUSSION

5.1. Utility–governance decoupling

Adoption is anchored in task value. Applicants appear willing to use AI recruitment when it connects them to relevant roles and can be tried or understood through experience. Accountability also matters, indicating that users value identifiable responsibility and recourse. Together, these effects describe conditional acceptance rather than unconditional enthusiasm.

5.2. Why fairness and trust may disappear in the full model

Fairness and trust may function as baseline expectations, indirect antecedents or post-outcome judgments rather than immediate adoption drivers. They may also share variance with accountability and trialability. A candidate can use a system because it is the only route to a desired job while still distrusting its decisions. Behavioural adoption should therefore not be mistaken for legitimacy.

5.4. Experience may reduce enthusiasm

Several descriptive contrasts show lower evaluations among respondents who report prior AI-recruitment experience. If replicated, this would indicate an experience–expectation gap: abstract AI promises appear attractive, but actual systems may feel opaque, rigid or difficult to contest. Longitudinal applicant-journey research is needed to test this possibility.

6. Implications

6.1. Implications for employers and vendors

Use AI only for job-relevant evidence and document why each feature predicts performance in the specific role.

Provide practice environments, sample assessments and accessible explanations before consequential evaluation.

Assign accountable human owners, establish appeal channels and record when human reviewers override automated recommendations.

Audit selection rates and errors across protected and inter-sectional groups; adoption volume is not evidence of fairness.

Minimise personal-data collection, state retention periods and separate optional enrichment data from information nec-

essary to assess the application.

6.2. Governance and policy implications

Employment AI requires lifecycle governance. The NIST functions—govern, map, measure and manage—offer an operational structure, while the EU AI Act demonstrates the regulatory significance of recruitment systems as high-risk AI. In India, consent, purpose limitation, data minimisation, security and applicant rights should be embedded in recruitment workflows under the Digital Personal Data Protection Act.

6.3. Theoretical implications

The findings qualify adoption theory by separating instrumental acceptance from normative legitimacy. Responsible-AI attributes may not all operate as direct predictors; fairness and transparency may influence adoption through accountability, perceived control or organisational trust. Future models should test those mechanisms instead of entering every ethical construct as a parallel predictor.

7. Conclusion and Future Scope

Applicants' adoption of AI recruitment is driven most strongly by job relevance, followed by trialability, complexity and accountability. Fairness, transparency, trust, anxiety and insecurity are non-significant in the full model. This pattern does not reduce the importance of responsible AI. It shows that people may adopt a useful system even when its legitimacy remains uncertain.

The study is limited by convenience sampling, heterogeneous respondent roles, self-report measures, one-time data collection, absent item provenance, unreported construct validity, no regression diagnostics and incomplete ANOVA reporting. The mixture of applicants and recruiters also blurs the focal decision perspective.

A Q1-level extension should separate actual applicants from recruiters, recruit participants immediately after a real selection encounter and validate the FAT-CAT and adoption constructs through CFA. A longitudinal design could test whether explanation and human review affect procedural justice, organisational trust, withdrawal and offer acceptance. Experimental audits should vary explanation, appeal, human involvement and model error while measuring subgroup fairness and actual candidate behaviour.

References

- Acikgoz, Y., Davison, K. H., Compagnone, M., & Laske, M. (2020). Justice perceptions of artificial intelligence in selection. *International Journal of Selection and Assessment*, 28(4), 399–416. <https://doi.org/10.1111/ijsa.12306>
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340. <https://doi.org/10.2307/249008>
- European Union. (2024). Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence. Official Journal of the European Union. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>
- Government of India. (2023). The Digital Personal Data Protection Act, 2023 (No. 22 of 2023). Ministry of Electronics and Information Technology. <https://www.meity.gov.in/content/digital-personal-data-protection-act-2023>
- Horodyski, P. (2023). Applicants' perception of artificial intelligence in the recruitment process. *Computers in Human Behavior Reports*, 11, Article 100303. <https://doi.org/10.1016/j.chbr.2023.100303>
- Köchling, A., & Wehner, M. C. (2020). Discriminated by an algorithm: A systematic review of discrimination and fairness by algorithmic decision-making in human resource management. *Business Research*, 13(3), 795–848. <https://doi.org/10.1007/s40685-020-00134-w>
- Langer, M., König, C. J., & Papatthanasious, M. (2019). Highly automated job interviews: Acceptance under the influence of stakes. *International Journal of Selection and Assessment*, 27(3), 217–234. <https://doi.org/10.1111/ijsa.12246>
- Lee, C. H., & Cha, K. J. (2023). FAT-CAT—Explainability and augmentation for an AI system: A case study on AI recruitment-system adoption. *International Journal of Human-Computer Studies*, 171, Article 102976. <https://doi.org/10.1016/j.ijhcs.2022.102976>
- Nagaraju, E., Kumari, T. L., Bambuwalla, S., & Rajalakshmi, M. (2024). AI-driven solutions for supply chain management. *Journal of Informatics Education and Research*, 4(2).
- National Institute of Standards and Technology. (2023). Artificial intelligence risk management framework (AI RMF 1.0) (NIST SP 100-1). U.S. Department of Commerce. <https://doi.org/10.6028/NIST.AI.100-1>
- Ore, O., & Sposato, M. (2022). Opportunities and risks of artificial intelligence in recruitment and selection. *International Journal of Organizational Analysis*, 30(6), 1771–1782. <https://doi.org/10.1108/IJOA-07-2020-2291>
- Pan, Y., Froese, F. J., Liu, G., Hu, Y., & Ye, M. (2022). The adoption of artificial intelligence in employee recruitment: The influence of contextual factors. *The International Journal of Human Resource Management*, 33(6), 1125–1147. <https://doi.org/10.1080/09585192.2021.1879206>
- Rigotti, C., & Fosch-Villaronga, E. (2024). Fairness, AI and recruitment. *Computer Law & Security Review*, 53, Article 105966. <https://doi.org/10.1016/j.clsr.2024.105966>
- Rogers, E. M. (2003). *Diffusion of innovations* (5th ed.). Free Press.
- van Esch, P., Black, J. S., & Arli, D. (2021). Job candidates' reactions to AI-enabled job application processes. *AI and Ethics*, 1(2), 119–130. <https://doi.org/10.1007/s43681-020-00025-0>
- Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS Quarterly*, 27(3), 425–478. <https://doi.org/10.2307/30036540>