



The Role of Artificial Intelligence System Reliability and Transparency in Reducing Algorithm Aversion and Enhancing Auditors' Trust in AI Outputs

Ridha Mohanad Al-Salman

Technical College of Management Kufa, Al-Furat Al-Awsat Technical University (ATU), Kufa, Iraq

* Corresponding Author: **Ridha Mohanad Al-Salman**

Article Info

ISSN (online): 2583-6641

Impact Factor (RSIF): 8.56

Volume: 05

Issue: 05

Received: 22-06-2026

Accepted: 24-07-2026

Published: 26-08-2026

Page No: 17-23

Abstract

Artificial intelligence (AI) is being used in audit processes. However, auditors may ignore algorithmic advice, a condition known as algorithm aversion. This study examines the association between the perceived dependability and transparency (explainability) of standard artificial intelligence systems and their effects on algorithm aversion and trust in AI outputs; and the association between trust and auditors' desire to rely professionally. The questionnaire was issued to 140 external and internal auditors and consisted of 25 items, measuring five constructs on a 5-point Likert scale. All constructs showed good internal consistency (Cronbach's $\alpha = 0.848-0.882$). Reliability and transparency were negatively related to algorithm aversion ($\beta = -0.29$ and -0.30) and positively related to trust ($\beta = 0.27$ and 0.31), while aversion was negatively related to trust ($\beta = -0.28$). Trust had the strongest standardized coefficient in the professional-reliance-intention model ($\beta = 0.40$). The models explained 24.7%, 45.6%, and 38.6% of the variance in aversion, trust, and reliance intentions, respectively. The percentile bootstrap analysis (5,000 resamples) revealed favorable indirect associations between reliability and transparency with reliance intention through trust. The confirmatory factor analysis (CFA) indicated a good fit ($\chi^2(265) = 283.74$, $p = 0.205$; CFI = 0.99; TLI = 0.99; RMSEA = 0.02). Composite reliability was ≥ 0.85 , average variance extracted was ≥ 0.53 , indicating a good convergent validity. The Fornell-Larcker criterion confirmed the discriminant validity. Analyses of diagnostics suggested sample adequacy (KMO = 0.90), no multicollinearity (maximum VIF = 1.84), and a Harman single factor of 37.3% for an initial assessment of common-method bias. All 8 hypotheses were empirically supported. Our results show that trust mediates the relationship between dependable and explainable AI design and auditors' intention to rely on AI outputs, but the cross-sectional design does not allow for causal conclusion.

DOI: <https://doi.org/10.54660/IJMOR.2026.5.5.17-23>

Keywords: artificial intelligence, auditing, algorithm aversion, trust in AI, Explainability, professional reliance intention, mediation

1. Introduction

Artificial intelligence is now a working part of modern auditing. Major accounting and audit firms are using machine learning, natural language processing, robotic process automation, and predictive analytics in order to extract data from documents, analyze entire populations of transactions, assess risk, and identify anomalies that could be missed by traditional sampling methods^[1, 3, 6]. Systematic and bibliometric studies suggest that artificial intelligence can improve operational efficiency, identify fraud better, improve the credibility of financial reporting, and change the auditor's function from retrospective inspection to continuous and real-time oversight^[2, 4]. Field experience shows that acceptance is inconsistent, with extensive use of "simple AI" such as optical character recognition and important data extraction while "complex AI" that supports judgment-intensive jobs is adopted more cautiously^[1].

The main barrier to realizing the potential of advanced AI in auditing is behavioural, not technological. Algorithm aversion^[11] is a phenomenon where decision makers tend to dismiss the advice from algorithms more than the advice from humans. Auditors are not immune: experimental studies show that auditors exposed to contradicting evidence from a firm's AI system rather than a human expert suggest less changes to management's complex estimates, a pattern that could harm audit quality^[10]. The audit process rests on professional skepticism and personal accountability, hence auditors frequently rely on their own judgement rather than only the output of the computer even if the system is correct^[5, 9]. Two design-related elements of the AI system are hypothesized to reduce this aversion. Reliability is the auditor's assurance that the system produces accurate results and has operated consistently. Transparency, or explainability, is the degree to which the system makes clear the technique used to get its findings. Research on trust in automation has indicated that both the perceived efficacy of a system and the clarity of its behavior are important aspects that affect acceptable reliance. The literature on explainable AI (XAI) states that opaque "black-box" models hinder trust, accountability and regulatory compliance in finance^[7, 8]. Studies conducted in the area of accounting demonstrate that the trustworthiness and clarity of a decision aid determine the willingness of an auditor to rely on it^[9]. However, despite the depth of theory, there is a paucity of studies that incorporate AI reliability and transparency, algorithm aversion, trust and professional reliance purpose within a single validated model of auditor behavior, and even fewer studies that investigate the indirect effect of trust. This study seeks to fill this gap using survey data from 140 internal and external auditors. The study aims to (i) assess the perception of AI reliability, transparency, algorithm aversion, trust and professional reliance intention; (ii) examine the influence of reliability and transparency on reduced aversion and increased trust; and (iii) examine the indirect effect of trust from system perceptions to reliance intention. The study provides a holistic and empirically supported methodology and practical guidance to audit firms, regulators, and system architects on how to match the design of AI tools to the professional psychology of auditors.

2. Materials and Methods

2.1. Theoretical Framework and Hypotheses Development

2.1.1. AI in accounting and auditing

The reviewers all agree on one thing: AI is transforming the evidential basis of auditing. Machine-learning algorithms tend to outperform traditional statistical techniques in prediction-oriented tasks like as fraud detection, going-concern assessments, and accounting estimates. They also enhance the quality of the accepted audit evidence through the exploitation of unstructured and textual data^[5]. AI in audit process reference architectures incorporate AI in the planning, risk assessment, control testing and substantive procedures with human supervision controls remaining in place^[3]. The literature on AI-enabled auditing has exploded since 2018 with studies addressing preparation and implementation, data-driven audit ecosystems, and issues of audit quality, professional skepticism, and ethical governance^[4]. In these evaluations Trust and perceived utility are consistently strong predictors of adoption. However, its adoption is accompanied by hazards of prejudice, deskilling,

and data privacy issues that require regulation^[2, 15, 16]. The rising competition between computer models and human predictions also pose institutional challenges to professional expertise and authority in the system of professions^[5, 17].

2.1.2. Algorithm aversion and trust in automation

The behavioural foundation of the study is built upon two literatures. Dietvorst, Simmons, and Massey^[11, 18] first showed that people lose confidence in an algorithm faster than in a human when seeing each err and thus shun superior algorithmic forecasters. Under different circumstances, the opposite phenomenon was found by Logg, Minson, and Moore^[12], called algorithm appreciation, showing that aversion is context-dependent and changeable. In auditing, Commerford *et al.*^[10] presented clear evidence of aversion, while Watson^[9] discussed how it may be overcome in accounting, identifying dependability and transparency as the two levers most likely to increase reliance. Second, the trust-in-automation literature shows that trust is mostly based on perceived system performance and process transparency, and calibrated trust is the one that leads to appropriate—neither excessive nor deficient—reliance^[13, 14]. Together, these literatures present dependability and transparency as antecedents of aversion and trust, and trust as the proximal correlate of reliance intention.

2.1.3. AI system reliability

Reliability is the confidence that the AI system is giving correct outputs and has a track record of performing well on similar jobs. Evidence of prior correctness substantially conditions confidence in an algorithm; disclosure of error rates and validated performance lessens hesitation to rely on it, and endorsement by a professional body or respectable firm improves perceived reliability further^[9, 11, 14]. Since aversion is mostly elicited by worries about whether the machine can be trusted to be correct, a system that is viewed as reliable should be associated with less aversion and greater trust:

H1: Auditors' algorithm aversion is inversely associated with perceived AI reliability.

H3: Auditors' faith in AI outputs is positively related to perceived AI reliability.

2.1.4. AI system transparency and explainability

Transparency measures the extent to which the system is able to explain its logic. In the XAI literature, opaque models are said to decrease trust, accountability and compliance, while explanations of the factors driving an outcome increase transparency and stakeholder trust^[7, 8]. In auditing, there is a particular value to explainability: If the result is explainable, the auditor can apply professional skepticism, critically analyze the evidence, and properly document it^[5, 9]. Hence, a transparent system should be related to lesser uncertainty and aversion, and better trust:

H2: Perceived AI transparency is inversely associated with auditor aversion to algorithms.

H4: There is a favorable relationship between auditors' trust in AI results and perceived AI transparency.

2.1.5. Algorithm aversion, trust, and professional reliance intention

Algorithm aversion as an inclination to distrust machine counsel should itself be related with decreased trust in AI outputs, system features held constant^[10, 11]. The key

correlate driving perceptions of the system to intended use is trust—the readiness of the auditor to treat AI outputs as reliable inputs to judgment, consistent with technology-acceptance theory and trust-in-automation models^[13, 15]. It is expected that trust leads to professional reliance intention, which is the extent to which the auditor intends to use AI outputs in the evaluation of evidence, not directly from system features. Thus, trust is expected to have the indirect relationships of reliability and transparency with reliance intention^[9, 14]. This reasoning leads to:

H5: Algorithm aversion is inversely related to auditors' trust in AI results.

H6: The trust in AI outcomes is favorably connected with auditors' professional reliance intention.

H7: Trust mediates positively the relationship between reliability and desire to rely professionally (REL → TRU → PRL).

H8: Transparency is positively indirectly associated with professional reliance intention through trust (TRA → TRU → PRL).

The structural model, therefore, postulates the following relationships: reliability and transparency as exogenous antecedents related to algorithm aversion (H1, H2) and trust (H3, H4); algorithm aversion as a correlate of trust (H5); trust as the proximal predictor of professional reliance intention (H6); and trust as the indirect pathway from reliability and transparency to reliance intention (H7, H8). Table 8 presents the estimated standardized coefficients for all eight pathways.

2.2. Research Methodology

Sample and design. The study adopted a quantitative survey research methodology. The population under study was professional auditors. In total, 140 valid responses were received from external auditors (55.7%) and internal auditors (44.3%) working in audit companies, banks and financial institutions, public sector, private corporations, self-employed audit offices and regulatory bodies. The sample met the conventional criteria for multiple regression analysis with a maximum of four predictors (minimum of 10-15 cases per predictor) and demonstrated excellent sampling adequacy (KMO = 0.90). Participation was voluntary and anonymous, respondents were informed about the academic purpose of the study and filling the questionnaire implied informed consent.

Constructions and Instruments. Perceptions were measured with a 25-item questionnaire using a five-point Likert scale ranging from 1 (strongly disagree) to 5 (strongly agree). Five latent constructs were measured using five questions each: the reliability of the AI system (REL), transparency and explainability of the AI system (TRA), algorithm aversion (AVR), trust in AI outputs (TRU) and intention to professionally rely on the AI system (PRL). Higher scores are associated with higher values of each component. The elements were conceptually based on research on algorithm aversion, auditor reliance, explainable artificial intelligence (XAI) and faith in automation^[8-14]. Table 1 shows one item for each build.

Table 1: Constructs and measurement model.

Code	Construct	Items	Representative item
REL	AI system reliability	5	"AI tools produce consistent results under similar audit conditions."
TRA	AI transparency & explainability	5	"The AI tool provides an understandable explanation of how it reached its result."
AVR	Algorithm aversion	5	"I prefer human auditor judgment even when the AI tool has higher historical accuracy."
TRU	Trust in AI outputs	5	"I regard AI outputs as dependable inputs when forming professional judgment."
PRL	Professional reliance intention	5	"I intend to use AI outputs in audit planning and risk assessment."

Review, Composite scores were derived by averaging the five elements that made up each construct. The data were examined using descriptive statistics, Cronbach's alpha, corrected item-total correlations, and Pearson correlations. Kaiser-Meyer-Olkin (KMO) measure and Bartlett's test of sphericity were used to test the appropriateness of sampling. Harman's single-factor result was not seen as convincing evidence of the absence of common-method bias but rather as an initial check. Confirmatory factor analysis (CFA) with maximum likelihood estimation was used to assess the measurement model. Standardized factor loadings, composite reliability (CR), average variance extracted (AVE), and overall fit indices (χ^2 /df, CFI, TLI, RMSEA). The Fornell-Larcker criterion and the heterotrait-monotrait ratio (HTMT) were used to assess the discriminant validity. The structural relationships were estimated using two complementary approaches: (a) three ordinary least squares models: Model A (aversion towards reliability and transparency), Model B (trust in reliability, transparency and aversion), Model C (intention to rely on reliability, transparency, aversion and trust); the variance inflation factors (VIF) were included for multicollinearity and HC3 heteroskedasticity-robust standard errors were used as a

sensitivity analysis; and (b) a latent variable structural equation model (SEM) that accounts for measurement error. Percentile-bootstrap confidence intervals, based on 5,000 resamples, were used to evaluate the two indirect pathways (H7, H8). All statistics were computed separately from the data at the item level

3. Results and Discussion

3.1. Results

3.1.1. Sample profile

The 140 respondents were experienced auditors. By role, 55.7% were external auditors and 44.3% internal auditors, drawn mainly from audit firms (36.4%), banks and financial institutions (15.7%), the public sector (14.3%), and private companies (14.3%). Most held a master's degree (42.1%) or bachelor's degree (32.9%), with 10.7% holding a doctorate. Experience was well distributed: 27.1% had fewer than five years, 32.1% five to nine years, 20.7% ten to fourteen years, and 20.0% fifteen years or more. AI familiarity was predominantly moderate (42.1%) or high (35.0%), and 45.7% had received AI-related training. Table 2 summarizes the profile.

Table 2: Sample profile of auditors (n = 140).

Audit type (%)	Qualification (%)	Experience (%)	AI familiarity (%)
External — 55.7	Master — 42.1	< 5 yrs — 27.1	Moderate — 42.1
Internal — 44.3	Bachelor — 32.9	5–9 yrs — 32.1	High — 35.0
—	PhD — 10.7	10–14 yrs — 20.7	Basic — 22.9
—	Higher dip./cert. — 14.3	≥ 15 yrs — 20.0	—

3.1.2. Reliability and descriptive statistics

All five categories had strong internal consistency: transparency ($\alpha = 0.882$), trust ($\alpha = 0.874$), reliability ($\alpha = 0.853$), algorithm aversion ($\alpha = 0.851$), and professional reliance intention ($\alpha = 0.848$). All corrected item-total correlations were >0.60 . The sampling adequacy was good ($KMO = 0.902$) and Bartlett's test was significant ($\chi^2(300) = 1849.29$, $p < 0.001$), indicating the applicability of the correlation matrix for factor analysis. The construct means

(Table 3) were 3.49, 3.43, 3.45 and 3.42 for dependability, transparency, trust and professional reliance intention respectively. Algorithm aversion was close to the neutral midpoint ($M = 2.90$) and did not significantly differ from 3.0, $t(139) = -1.24$, $p = 0.218$, and should therefore be considered approximately neutral, not clearly low. All composite-score distributions demonstrated minor skewness and kurtosis ($|\text{skewness}| \leq 0.34$; $|\text{kurtosis}| \leq 0.83$).

Table 3. Construct reliability and descriptive statistics (n = 140).

Construct	α	Mean	SD	Skew.	Kurt.
AI reliability (REL)	0.853	3.491	0.932	-0.15	-0.83
AI transparency (TRA)	0.882	3.427	0.997	-0.28	-0.80
Algorithm aversion (AVR)	0.851	2.899	0.969	0.04	-0.81
Trust in AI outputs (TRU)	0.874	3.446	0.983	-0.33	-0.56
Professional reliance intention (PRL)	0.848	3.420	0.922	-0.21	-0.47

3.1.3. Correlations and discriminant validity

All constructs were highly associated in the predicted ways (Table 4). Reliability and transparency were positively associated to trust ($r = 0.516$ and 0.540) and reliance intention ($r = 0.451$ and 0.381) and adversely related to aversion ($r = -0.414$ and -0.422). Trust was favorably connected with reliance intention ($r = 0.582$) and negatively correlated with

aversion ($r = -0.526$), and aversion was negatively correlated with reliance intention ($r = -0.440$). All HTMT ratios were less than or equal to 0.675 (Table 6), indicating early evidence for discriminant validity. Harman's single component accounted for 37.3% of the variance. This diagnostic cannot rule out common-method bias by itself, and the finding should be viewed with caution.

Table 4: Pearson correlation matrix (n = 140).

	REL	TRA	AVR	TRU	PRL
REL	1.000	0.415	-0.414	0.516	0.451
TRA	0.415	1.000	-0.422	0.540	0.381
AVR	-0.414	-0.422	1.000	-0.526	-0.440
TRU	0.516	0.540	-0.526	1.000	0.582
PRL	0.451	0.381	-0.440	0.582	1.000

3.1.4. Regression models and hypothesis tests

Table 5 shows the three regression models. Variance Inflation Factors (VIFs) for all variables were below 1.84 showing the lack of multicollinearity. In Model A, dependability ($\beta = -0.288$, $p < 0.001$) and transparency ($\beta = -0.302$, $p < 0.001$) significantly predicted less algorithm aversion, accounting for 24.7% of the variance. In Model B, trust was positively related to dependability ($\beta = 0.271$, $p < 0.001$) and transparency ($\beta = 0.308$, $p < 0.001$) and negatively related to aversion ($\beta = -0.283$, $p < 0.001$), accounting for 45.6% of the variance in trust. In Model C, trust showed the highest

standardized coefficient for intention of professional reliance ($\beta = 0.400$, $p < 0.001$), reliability showed a lesser correlation ($\beta = 0.171$, conventional $p = 0.038$; HC3 $p = 0.049$), and the direct effects of transparency ($\beta = 0.033$, $p = 0.694$) and aversion ($\beta = -0.146$, $p = 0.077$) were not significant after the inclusion of trust, demonstrating a pattern of indirect transmission through trust. The model explained 38.6% of the variance in reliance intention. The results for H1-H6 remained the same when HC3-robust standard errors were used. The assumptions H1-H6 were confirmed (Table 8).

Table 5. Regression results (standardized β ; n = 140).

Predictor	Model A DV: Aversion	Model B DV: Trust	Model C DV: Reliance int.
Reliability (REL)	-0.288***	0.271***	0.171*
Transparency (TRA)	-0.302***	0.308***	0.033
Algorithm aversion (AVR)	—	-0.283***	-0.146
Trust (TRU)	—	—	0.400***
R ²	0.247	0.456	0.386
F	22.44***	37.93***	21.26***

Note: * $p < 0.05$, *** $p < 0.001$. "—" indicates the predictor was not entered in that model. HC3 heteroskedasticity-robust standard errors yielded the same significance decisions for H1-H6.

Indirect associations. Bootstrap analysis with 5,000 resamples supported both indirect paths (Table 7). The indirect association of reliability with reliance intention through trust was positive (0.106, 95% CI [0.041, 0.18⁷¹]), as was that of transparency (0.114, 95% CI [0.047, 0.19⁶¹]); the indirect association of aversion with reliance intention through trust was negative (−0.109, 95% CI [−0.200, −0.04²¹]). Because transparency's direct effect on reliance

intention was non-significant while its indirect association was supported, trust accounts for the full transparency–reliance-intention association, whereas it accounts for part of the reliability association. Hypotheses H7 and H8 were supported. These indirect patterns are consistent with a mediating role for trust but, given the cross-sectional design, are interpreted as statistical rather than causal pathways.

Table 6: Discriminant validity (HTMT ratios).

	REL	TRA	AVR	TRU	PRL
REL	—				
TRA	0.476	—			
AVR	0.484	0.488	—		
TRU	0.599	0.615	0.609	—	
PRL	0.530	0.438	0.517	0.675	—

Table 7: Bootstrap indirect associations (5,000 resamples).

Indirect path	Effect	95% CI	Via trust
REL → TRU → PRL	0.106	[0.041, 0.18 ⁷¹]	Partial
TRA → TRU → PRL	0.114	[0.047, 0.19 ⁶¹]	Full
AVR → TRU → PRL	−0.109	[−0.200, −0.04 ²¹]	Present

Table 8: Summary of hypothesis tests.

H	Path	Expected	β / indirect	Decision
H1	Reliability → Algorithm aversion	Negative	−0.288***	Supported
H2	Transparency → Algorithm aversion	Negative	−0.302***	Supported
H3	Reliability → Trust	Positive	0.271***	Supported
H4	Transparency → Trust	Positive	0.308***	Supported
H5	Algorithm aversion → Trust	Negative	−0.283***	Supported
H6	Trust → Reliance intention	Positive	0.400***	Supported
H7	Reliability → Trust → Reliance int.	Positive (ind.)	0.106*	Supported
H8	Transparency → Trust → Reliance int.	Positive (ind.)	0.114*	Supported

Note: *** p < 0.001; * bootstrap 95% CI excludes zero.

3.1.5. Confirmatory factor analysis and structural validation

A confirmatory factor analysis of the five-factor measurement model demonstrated an excellent fit to the data: $\chi^2(265) = 283.74$, $p = 0.205$ (non-significant), CFI = 0.989, TLI = 0.988, and RMSEA = 0.023. All standardized loadings were significant and ranged from 0.65 to 0.86 ($p < 0.001$). Convergent validity was established, with composite

reliability (CR) between 0.851 and 0.883 (all > 0.70) and average variance extracted (AVE) between 0.533 and 0.602 (all > 0.50); Table 9 reports these indices. Discriminant validity was confirmed by the Fornell–Larcker criterion (Table 10): the square root of each construct's AVE exceeded its correlations with every other construct, complementing the HTMT evidence in Table 6.

Table 9: Confirmatory factor analysis: convergent validity (n = 140).

Construct	Loadings	CR	AVE	√AVE
AI reliability (REL)	0.65–0.79	0.853	0.539	0.734
AI transparency (TRA)	0.69–0.86	0.883	0.602	0.776
Algorithm aversion (AVR)	0.70–0.77	0.851	0.533	0.730
Trust in AI outputs (TRU)	0.70–0.84	0.874	0.583	0.763
Professional reliance intention (PRL)	0.68–0.78	0.851	0.533	0.730

Note. CR = composite reliability; AVE = average variance extracted. All standardized loadings significant at $p < 0.001$.

Table 10: Discriminant validity — Fornell–Larcker criterion.

	REL	TRA	AVR	TRU	PRL
REL	0.734				
TRA	0.415	0.776			
AVR	−0.414	−0.422	0.730		
TRU	0.516	0.540	−0.526	0.763	
PRL	0.451	0.381	−0.440	0.582	0.730

Note. Diagonal (bold) = square root of AVE; off-diagonal = inter-construct correlations. Discriminant validity holds where $\sqrt{\text{AVE}}$ exceeds the correlations in its row and column.

The complete model was re-estimated as a latent-variable Structural Equation Model (SEM) to confirm the regression results with correction for measurement error. The SEM

reproduced the pattern of results with the same excellent fit (CFI = 0.989, TLI = 0.988, RMSEA = 0.023): all six direct paths were significant in the expected directions (REL → AVR β = -0.34; TRA → AVR β = -0.31; REL → TRU β = 0.28; TRA → TRU β = 0.33; AVR → TRU β = -0.32; TRU → PRL β = 0.49, all p < 0.005), and the direct latent effects of reliability and transparency on reliance intention were non-significant once trust was added. This result further supports the interpretation that trust mediates the influence of reliability and transparency on reliance intention. The convergence of estimates from the ordinary least squares, bootstrap, and latent SEM models increases confidence in the robustness of the model.

3.2. Discussion

The findings corroborate all eight assumptions and indicate a consistent pattern in which trust is the key relationship between the perceived design of an AI system and its intended professional application. As expected in H1 and H2, both reliability and transparency were related to less algorithm aversion. This is consistent with experimental and accounting-specific evidence that the most effective way to reduce auditors' reluctance to accept machine advice is through accuracy and an intelligible process [9, 10, 11]. As expected from H3-H5, reliability and transparency were positively related to trust, whereas aversion was negatively related to it. This is in line with the trust-in-automation view that performance and process transparency are the primary bases of trust [13, 14], and with the XAI argument that explainability is an essential component of trust in financial AI [7, 8].

The indirect results are the main theoretical contribution of the paper. Trust had the highest coefficient for professional reliance intention (H6). The bootstrap analysis suggested that dependability and transparency are substantially connected to reliance intention through trust (H7, H8). The pattern in Model C is particularly instructive: trust was not significantly directly associated with reliance intention after controlling for transparency (trust accounts for the complete transparency relationship), but dependability kept a small direct association on top of its indirect path. In practice, explainability of an AI system does not, in itself, increase auditors' intention to rely on it; explainability seems to work by establishing confidence, and it is trust that connects capable, transparent systems to intended use. This indicates that the key, and potentially fragile, ingredient is trust in the chain from AI design to audit practice, however the cross-sectional design means these are statistical rather than obviously causative routes.

This suggests plausible levers for reducing algorithm aversion for audit firms and AI vendors: investments in demonstrable accuracy (validation reports, performance histories, professional endorsement), and in genuine explainability (intelligible drivers of each output), but especially insofar as these translate into calibrated trust, and always in conjunction with the human-in-the-loop review that professional standards require [3, 6, 14]. This conclusion, that these auditors were on average around neutral in aversion but nonetheless conditioned their reliance intention on trust, indicating that the profession is receptive to AI when it is seen as trustworthy and transparent, rather being reflexively opposed to it.

4. Conclusions

This study modeled the association between perceived reliability and transparency of standard artificial intelligence systems and auditors' algorithm aversion, trust, and professional reliance intention and if trust mediates the association between design and reliance. Based on responses from 140 external and internal auditors, it found that reliability and transparency were both associated with less algorithm aversion and more trust, that aversion was associated with less trust, and that trust was the strongest correlate of professional reliance intention, accounting for the full transparency association and part of the reliability association with reliance intention. The measurement model was reliable and demonstrated preliminary validity (α = 0.848–0.882; KMO = 0.90; HTMT ≤ 0.68; VIF < 1.84) and all eight assumptions were supported. The key addition is the proof that trust is the mechanism that links trustworthily explainable AI design to auditors' intention to depend on AI. "The results are clearly significant. Audit businesses would do well to look for AI tools that log their accuracy and explain their thinking, and supplement their use with training that helps auditors assess, rather than reflexively dismiss or overtrust, machine advice. Regulators and standard-setters should foster transparent, auditable and verified systems, in line with professional skepticism, and accounting education can equip future auditors to evaluate and manage AI outputs. These conclusions are qualified by several restrictions, and suggest future studies. The design is cross-sectional, which prevents causal inference, and is based on self-reported perceptions and intents that may not reflect real reliance behaviour; experimental and longitudinal studies would improve causal assertions. Harman's test and HTMT are only preliminary evidence of the common-method bias and discriminant validity, and some items need further refinement (e.g., two trust items refer to reliability and explainability, and may partly overlap with those antecedents; some aversion items may tap legitimate professional skepticism rather than aversion per se). Future instruments should purify these items and undergo expert content validation (CVI/CVR) and back-translation. Finally, the sample is auditor-specific and concentrated in one national setting. Larger, multi-country samples analysed with confirmatory structural-equation modelling (e.g., PLS-SEM with Q^2 and PLSpredict) would allow a fuller test of the model and its moderators, such as task type and auditor experience.

Declarations

Ethics and consent. Participation was voluntary and anonymous; respondents were informed of the study's academic purpose, and completion of the questionnaire constituted informed consent. No personally identifying information was collected.

Conflict of interest. The author declares no conflict of interest.

Funding. This research received no specific grant from any funding agency in the public, commercial, or not-for-profit sectors.

Data availability. The anonymized dataset and analysis outputs supporting the findings are available from the author upon reasonable request.

References

1. Kokina J, Blanchette S, Davenport TH, Pachamanova D. Challenges and opportunities for artificial intelligence in auditing: Evidence from the field. *Int J Account Inf Syst.* 2025.
2. Abdo-Salloum M, Chehade S. The role of artificial intelligence in transforming accounting and auditing practices: A systematic review. *SAGE Open.* 2026;16(1). doi:10.1177/21582440251403296.
3. Anwar, Akeel MO. Integrating artificial intelligence in audit workflow: Opportunities, architecture, and challenges: A systematic review. *Account Audit.* 2026;2(1):4.
4. Sundarasan S, Kamaludin K, Nakiran D. From adoption to audit quality: Mapping the intellectual structure of artificial intelligence-enabled auditing. *J Risk Financ Manag.* 2026;19:209.
5. Jones S, Free C. What accountants need to know about artificial intelligence and machine learning: A review and call for future research. *J Account Lit.* 2026.
6. Nartey OL, Aryeetey S, Apaflo DA, Yeboah MM. AI and blockchain in financial auditing: A systematic review of fraud detection techniques and their impact on investor confidence in the U.S. market. *EPRA Int J Res Dev.* 2026;11(3). doi:10.36713/epra26771.
7. Tuge NB, Msweli NT. Explainable artificial intelligence in the banking sector: A systematic literature review. *Appl Artif Intell.* 2026;40(1):e2676353. doi:10.1080/08839514.2026.2676353.
8. Zhang C, Cho S, Vasarhelyi M. Explainable artificial intelligence (XAI) in auditing. *Int J Account Inf Syst.* 2022;46:100572. doi:10.1016/j.accinf.2022.100572.
9. Alhadrawi AKA, et al. Unveiling extremism: Leveraging digital data mining strategies. *J Ecohumanism.* 2024;3(7):492-502.
10. Watson DE. Through the looking glass: Overcoming algorithm aversion in accounting [doctoral dissertation]. Tampa (FL): University of South Florida; 2024.
11. Commerford BP, Dennis SA, Joe JR, Ulla JW. Man versus machine: Complex estimates and auditor reliance on artificial intelligence. *J Account Res.* 2022;60(1):171-201. doi:10.1111/1475-679X.12407.
12. Dietvorst BJ, Simmons JP, Massey C. Algorithm aversion: People erroneously avoid algorithms after seeing them err. *J Exp Psychol Gen.* 2015;144(1):114-126. doi:10.1037/xge0000033.
13. Logg JM, Minson JA, Moore DA. Algorithm appreciation: People prefer algorithmic to human judgment. *Organ Behav Hum Decis Process.* 2019;151:90-103. doi:10.1016/j.obhdp.2018.12.005.
14. Lee JD, See KA. Trust in automation: Designing for appropriate reliance. *Hum Factors.* 2004;46(1):50-80. doi:10.1518/hfes.46.1.50_30392.
15. Hoff KA, Bashir M. Trust in automation: Integrating empirical evidence on factors that influence trust. *Hum Factors.* 2015;57(3):407-434. doi:10.1177/0018720814547570.
16. Nasser QO, Al-Hadrawi BK. Thought leadership and its impact on reducing customer incivility: An analytical study in Al-Qasim General Hospital/Iraq. In: *AIP Conference Proceedings.* Melville (NY): AIP Publishing; 2024. p. 080004.
17. Davis FD. Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Q.* 1989;13(3):319-340.
18. Farmer RL, Lockwood AB, Goforth A, Thomas C. Artificial intelligence in practice: Opportunities, challenges, and ethical considerations [preprint]. 2024.
19. Abbott A. The system of professions: An essay on the division of expert labor. Chicago (IL): University of Chicago Press; 1988.

How to Cite This Article

Al-Salman RM. The role of artificial intelligence system reliability and transparency in reducing algorithm aversion and enhancing auditors' trust in AI outputs. *International Journal of Management and Organizational Research.* 2026;5(5):17-23. doi:10.54660/IJMOR.2026.5.5.17-23.

Creative Commons (CC) License

This is an open access journal, and articles are distributed under the terms of the Creative Commons Attribution-NonCommercial-ShareAlike 4.0 International (CC BY-NC-SA 4.0) License, which allows others to remix, tweak, and build upon the work non-commercially, as long as appropriate credit is given and the new creations are licensed under the identical terms.