

Artificial Intelligence in Financial Reporting Fraud Detection: An Empirical Investigation of Auditor Trust as a Mediating Mechanism in U.S. Markets

**Mohammad Musa Mia¹, Md Zahurul Islam², Abu Sayed³, Faysal Ahmed⁴,
Ataur Rahman⁵, Mohammad Ariful Aziz⁶ and Md Samirul Islam⁶**

Abstract

Artificial intelligence (AI) is increasingly used to support financial reporting fraud detection, yet algorithmic accuracy alone does not determine how effective AI proves to be in practice. This study develops and tests an integrated model in which Auditor Trust in AI Systems mediates the effects of AI Predictive Capability, Explainable AI, Data Quality, and AI Governance and Internal Control Quality on Perceived Financial Reporting Fraud Detection Effectiveness. The model draws on human–AI trust theory, the automation–augmentation perspective, and agency theory. Data from 450 U.S. auditing and financial reporting professionals were collected by structured questionnaire and analysed using Partial Least Squares Structural Equation Modelling in SmartPLS 4. All four antecedents significantly predicted Auditor Trust in AI Systems ($R^2 = .697$), which in turn strongly predicted fraud detection effectiveness ($R^2 = .524$; $\beta = .724$, $p < .001$) and significantly mediated the effects of all four antecedents. The findings show that AI-enabled fraud detection is a socio-technical outcome that delivers value only when systems are explainable, built on sound data, embedded in appropriate governance structures and trusted by their users. The study contributes an integrated framework linking technological, informational, organisational and behavioural conditions, with practical guidance for audit firms and regulators.

JEL classification numbers: M41; M42; O33; G30.

Keywords: artificial intelligence, financial reporting fraud detection, auditor trust, explainable AI, data quality, AI governance.

¹ School of Business, International American University.

² Accounting, University of New Haven.

³ Doctorate of Business Administration, International American University.

⁴ School of Accounting, Finance, Economics, & Decision Sciences, Western Illinois University.

⁵ School of Business, International American University.

⁶ Department of Management Information Systems, International American University.

1. Introduction

Given the important role of financial reporting in reducing the information asymmetry that underpins investor confidence in capital markets (Jensen and Meckling, 1976; Dechow et al., 2011), deception in financial reporting remains a notable threat: it distorts corporate performance, misleads stakeholders and jeopardises market integrity (Perols, 2011; Bao et al., 2020). In the American market, where companies must satisfy elaborate reporting and regulatory requirements (Dechow et al., 2011), quick and reliable fraud detection has become increasingly critical for auditors, regulators and investors.

Traditional fraud detection approaches rely heavily on auditor judgment, internal control mechanisms and statistical methods; however, these techniques are often challenged by vast and complex financial data sets (Perols, 2011; Cecchini et al., 2010). Artificial intelligence offers a better solution that uses machine learning, anomaly detection and predictive analytics to identify unusual reporting patterns and potential errors (Bao et al., 2020). Existing studies suggest that AI-based models can improve the precision of detection and support audit decision-making with greater efficiency (Cecchini et al., 2010). However, AI in fraud detection goes beyond technology. It depends on factors such as explainability, data integrity, governance and user trust (Glikson and Woolley, 2020; Zhou et al., 2023). According to Zhou et al. (2023), auditors may be less assured when evaluating fraud indicators produced by opaque systems whose reasoning they do not understand. Poor data quality or weak governance may likewise inhibit the trustworthiness of results obtained from AI (Wang and Strong, 1996; Tan et al., 2023).

While the literature on AI-enabled fraud detection is now relatively abundant, three constraints remain. First, past research has predominantly evaluated AI and machine-learning systems using technical orientations: classification accuracy, effectiveness of predictions and comparison with existing statistical methods (Perols, 2011; Bao et al., 2020). As a consequence, the conditions under which technically advanced systems deliver effective and professionally applied fraud-detection outputs have been studied only in a contested manner (Glikson and Woolley, 2020; Raisch and Krakowski, 2021). Second, technological, informational and organisational factors in AI adoption have mostly developed as separate research streams, leading to very limited integration of AI capability with explainability, data quality and governance within a single explanatory framework (Wang and Strong, 1996; Arrieta et al., 2020). Third, the link between these antecedents and fraud-detection outcomes remains loosely established: prior research indicates that system attributes can affect trust, but offers little empirical evidence of trust operating as a mediating construct (Glikson and Woolley, 2020; Raisch and Krakowski, 2021). Relatedly, Auditor Trust in AI Systems has seldom been investigated empirically using data drawn directly from experienced practitioners.

This research contributes to filling these gaps with a survey-based study of qualified auditing and financial reporting professionals in the U.S. market, developing and

testing an integrated framework in which AI Predictive Capability, Explainable AI, Data Quality, and AI Governance and Internal Control Quality affect Perceived Financial Reporting Fraud Detection Effectiveness through Auditor Trust in AI Systems. The study goes beyond the prevalent frame of technical performance by examining trust as the mediator through which technological, informational and organisational conditions translate into professional reliance and perceived effectiveness. The remainder of the paper reviews the literature and develops the hypotheses in Section 2; describes the methodology in Section 3; presents the empirical results in Section 4; discusses the findings in Section 5; and closes in Section 6 with theoretical and managerial implications, limitations and directions for future research.

2. Literature Review

2.1 AI-Enabled Financial Reporting Fraud Detection

AI is increasingly redefining the landscape of financial reporting, audit and fraud detection, since it can analyse large sets of financial data and identify patterns that cannot easily be discovered through conventional analytical methods (Vasarhelyi et al., 2015; Bao et al., 2020). Firms engaging in deceptive behaviour may purposely hide inconsistencies within complex transactions, obliging auditors and regulators to review vast datasets under time and resource constraints (Perols, 2011; Bao et al., 2020). The rapid growth and intricacy of financial data (Cecchini et al., 2010; Vasarhelyi et al., 2015) increasingly lead to the failure of traditional approaches — auditor judgment, corporate governance review, accounting ratios and historical analysis — in detecting complex, nonlinear fraud patterns. This has generated interest in machine-learning approaches that are effective in anomaly detection, classification and predictive modelling (Perols, 2011; Bao et al., 2020). This body of literature demonstrates the forecasting potential of AI but remains fragmented across computational fraud detection, explainable AI, data integrity, governance and human confidence in intelligent systems (Wang and Strong, 1996; Arrieta et al., 2020; Glikson and Woolley, 2020), and therefore lacks substantial investigation of the conditions under which practitioners would generate trust in, and use, AI-generated fraud indicators (Bao et al., 2020; Raisch and Krakowski, 2021).

A growing body of recent applied research corroborates this viewpoint from multiple vantage points, but tends to be more descriptive and technical than explanatory. Empirical studies show that AI and machine-learning technologies help detect anomalies in financial data (Antwi et al., 2024; Dash et al., 2025) by identifying discrepancies and unusual transactions and by supporting fraud detection, risk assessment and audit quality (Raza et al., 2025). Research has demonstrated that algorithmic approaches to the analysis of financial statements extend traditional ground (Chen et al., 2025; De Araújo, 2025), that artificial intelligence is being incorporated into audit processes through AI-assisted auditing tools (Gu et al., 2024), that AI implementation is altering the quality of financial reporting (Oyeniyi et al., 2024) and that it is disrupting financial auditing

frameworks (Priyo, 2025; Najm, 2025). This literature is largely silent about the non-technical factors — explainability, data quality, governance and trust — that turn technical capability into an actual facilitator of professional fraud detection in practice, which is the gap this study addresses.

2.2 AI Predictive Capability

AI Predictive Capability refers to the ability of AI-based systems to analyse complex transaction data, detect unusual behaviour and generate immediate, actionable fraud-risk information (Cecchini et al., 2010; Bao et al., 2020). This functionality relies on the ability of machine-learning models to evaluate nonlinear relationships among multiple financial conditions that are difficult to analyse with conventional analytical tools (Perols, 2011; Bao et al., 2020). Earlier research suggests that, under certain empirical conditions, machine-learning-based methods can classify fraudulent firms and improve fraud-risk prediction (Cecchini et al., 2010; Bao et al., 2020). When users believe a technology is competent for its purpose — a function of their perception of its ability and anticipated dependability — they also view its output as trustworthy in professional evaluations (Glikson and Woolley, 2020). This conceptual structure aligns with the broader organisational literature on AI capabilities, in which perceived analytical skill is an independent precursor to downstream performance outcomes rather than a substitute for them (Mikalef and Gupta, 2021), and it is supported by emerging evidence showing that algorithmic techniques in financial-statement analysis can identify violations of economic patterns that traditional methods would likely overlook (Chen et al., 2025; De Araújo, 2025).

Predictive ability needs to be conceptually separated from fraud-detection effectiveness: the former concerns the quality of the analysis, while the latter concerns how useful AI-generated suggestions are in real applications. Computational analyses provide evidence of the predictive capacity of AI but offer little evidence on the role of professional perception of that capacity. This study therefore examines AI Predictive Capability as an antecedent of auditor trust in AI systems:

H1. *Auditors place greater trust in AI systems that demonstrate stronger predictive ability.*

2.3 Explainable AI

Explainable artificial intelligence (XAI) relates to the extent to which users are able to understand, inspect and evaluate the rationale, procedures and outputs of AI systems (Arrieta et al., 2020). Explainability grows in importance as complex models produce accurate predictions while providing limited insight into how those outputs are generated, raising concerns about transparency, responsibility and over-reliance in important decision-making settings (Arrieta et al., 2020; Glikson and Woolley, 2020). This matters particularly in fraud detection, where AI-generated

risk evaluations influence audit strategies, choices of investigation and the handling of potentially fraudulent disclosures; practitioners may balk at algorithmic recommendations that lie beyond their understanding (Arrieta et al., 2020). In this context, explainability may increase perceived clarity and reduce the uncertainty associated with algorithmic recommendations (Glikson and Woolley, 2020).

The relationship is not a direct equation: overly complex explanations may increase cognitive load, while overly facile explanations may create inaccurate views of competence (Arrieta et al., 2020; Bauer et al., 2023). In general, however, the literature reports a positive association with trust. A recent meta-analysis of experimental and field results shows a consistent positive association between explainability and user trust in AI systems (Atf and Lewis, 2025), and complementary research on warranted trust shows that explanations can calibrate reliance rather than simply increase it (Duarte et al., 2022). In accounting and auditing, explainability has been studied as a tool to foster trust and regulatory compliance (Hatahali, 2025; Vardhan, 2021), as a point of distinction between competing XAI-based methods for financial-audit tasks (Ganapathy, 2025), as a mediator of auditor judgment quality (Nugrahanti et al., 2025), and as a lever for transparency and trust in fraud detection in the banking industry (S et al., 2025). This body of evidence supports:

H2. *Explainable AI increases auditor trust in AI systems.*

2.4 Data Quality

Data Quality refers to the extent to which data are accurate, complete, consistent, relevant, timely and fit for their intended use (Wang and Strong, 1996). The effectiveness of AI-based fraud detection relies heavily on the quality of the data used to train, validate and run analytical models, because even the most sophisticated algorithms produce misleading output from invalid input (Wang and Strong, 1996; Vasarhelyi et al., 2015). An increase in data volume alone does not mean that better evidence and improved decision-making will follow (Vasarhelyi et al., 2015), and false, incomplete, conflicting or out-of-date data reduce the trustworthiness of AI-generated fraud indicators and, subsequently, professional trust in audit findings. Data Quality is therefore evaluated not only as a preprocessing issue but as an information state that affects perceived trustworthiness. This perspective is consistent with recent systematic reviews of data-analytics use in financial-statement fraud detection, which identify data quality as a recurring methodological weakness that undermines models (Gkegkas et al., 2025), and with work on AI-assisted decision support that links data quality to trust and provenance in complex decision environments (Bondac et al., 2026). Consequently:

H3. *Data quality increases auditor trust in AI systems.*

2.5 AI Governance and Internal Control Quality

AI governance refers to the rules, principles, processes and accountability mechanisms by which organisations manage the development and use of AI systems, confronting questions of decision-making rights, transparency, human oversight and accountability for algorithmic outcomes (Shrestha et al., 2019). Systems of internal control are intended to provide complementary protections for risk management and accountability; while AI does not remove the need for controls, it introduces additional requirements around data integrity, system access, model oversight and assessment of algorithmic output. Strong governance and robust internal controls define roles and responsibilities, enable human intervention (Shrestha et al., 2019; Raisch and Krakowski, 2021) and reduce ambiguity by clarifying who is permitted to use AI, whereas weak governance contributes to uncertainty surrounding data integrity, accountability and liability for errors made in AI-based decisions.

Recent research on the internal audit function indicates that AI is not simply a new technology offering organisations unprecedented opportunities, but also a source of governance challenges that must be addressed proactively rather than neglected (Ahmed and Abdellah, 2025), and empirical evidence shows that strengthening internal control systems alongside AI implementation improves accounting information systems (Hastuti and Widuri, 2025). Extensive surveys of audit practice show that the challenges and opportunities associated with AI adoption are closely intertwined with firm-level governance decisions (Kokina et al., 2025), that AI is reshaping how audit firms organise internally and supervise the work of senior and junior staff (Law and Shen, 2025), and that the quality of AI-supported audits depends on the governance frameworks in place and on the characteristics of individual auditors (Khosravi et al., 2025). Related work raises concerns about ethical and regulatory oversight (Imane, 2025), which emphasises the importance of high-quality governance and control as an independent organisational antecedent. This study proposes the integration of these interrelated organisational elements into a single construct:

H4. *The robustness of AI governance and internal control processes increases auditor trust in AI systems.*

2.6 Auditor Trust in AI Systems and Fraud Detection Effectiveness

Trust plays a central role in human–AI research because people must decide how much faith to place in intelligent systems under conditions of uncertainty and limited knowledge of the systems' inner mechanisms (Glikson and Woolley, 2020), particularly when AI-generated results affect high-stakes decisions that cannot readily be verified. Auditor Trust in AI Systems (ATAIS) is defined as the willingness of auditors to rely on outputs generated by AI, and it is driven by auditors' perceptions of the competence, reliability and predictability of AI systems and by perceptions of appropriate system behaviour. In the human–AI trust

literature, key antecedents include perceived accuracy, reliability, transparency and predictability (Glikson and Woolley, 2020), consistent with the widely used technology acceptance model in which perceived usefulness and ease of use are significant determinants of professionals' willingness to rely on new information technologies (Davis, 1989). Recent survey results suggest that auditors' perceptions of AI are closely related to professional judgment (Sahla and Latif, 2023) and that higher fraud-detection performance is obtained when data analytics and auditor expertise are combined rather than substituted for one another (Suyono et al., 2025). Financial Reporting Fraud Detection Effectiveness refers to the extent to which fraud-detection mechanisms identify potentially fraudulent reporting accurately, efficiently and in a timely manner, enabling further investigation and decision-making by the relevant parties (Perols, 2011; Bao et al., 2020). To use these signals, professionals must trust, interpret and integrate them into expert judgment; trust is therefore considered a key factor influencing perceived effectiveness.

H5. *Auditor trust in AI systems increases the perceived effectiveness of financial reporting fraud detection.*

2.7 Mediating Role of Auditor Trust in AI Systems

The proposed model identifies Auditor Trust in AI Systems as the route through which technological, informational and organisational factors affect fraud-detection effectiveness. AI predictive capability may enhance beliefs about system efficacy, explainable AI may reduce ambiguity and support professional assessment, data quality may provide assurance of the reliability of analytical results, and AI governance and internal control quality may provide organisational assurance of accountability and responsible system use (Wang and Strong, 1996; Shrestha et al., 2019; Arrieta et al., 2020; Glikson and Woolley, 2020).

In these circumstances, favourable system attributes do not necessarily improve fraud-detection outcomes on their own, since professionals still need to determine how reliable AI-generated results are before acting on them. Trust therefore serves as a theoretically robust mediating mechanism between the antecedents of AI-enabled fraud detection and perceptions of professional outcomes.

H6a. *Auditor trust in AI systems mediates the relationship between AI predictive capability and the perceived effectiveness of financial reporting fraud detection.*

H6b. *Auditor trust in AI systems mediates the relationship between explainable AI and the perceived effectiveness of financial reporting fraud detection.*

H6c. *Auditor trust in AI systems mediates the relationship between data quality and the perceived effectiveness of financial reporting fraud detection.*

H6d. *Auditor trust in AI systems mediates the relationship between AI governance and internal control quality and the perceived effectiveness of financial reporting fraud detection.*

2.8 Theoretical Integration

The model derives from the combined influence of three theoretical perspectives. Agency theory explains that information asymmetry and the separation of ownership and control create scope for managerial opportunism and generate a need for strong monitoring and fraud detection in a reporting context (Jensen and Meckling, 1976). The automation–augmentation perspective explains that AI creates organisational value through the combination of technical capability and expert judgment rather than through automation alone (Raisch and Krakowski, 2021). The human–AI trust framework explains the behavioural processes through which technological, informational and organisational factors translate into professional reliance and perceived effectiveness (Glikson and Woolley, 2020). Together, these perspectives support a model in which AI Predictive Capability and Explainable AI represent technical features, Data Quality represents the reliability of informational resources, and AI Governance and Internal Control Quality represents organisational conditions, all of which affect Auditor Trust in AI Systems and, through it, Perceived Financial Reporting Fraud Detection Effectiveness (Figure 1).

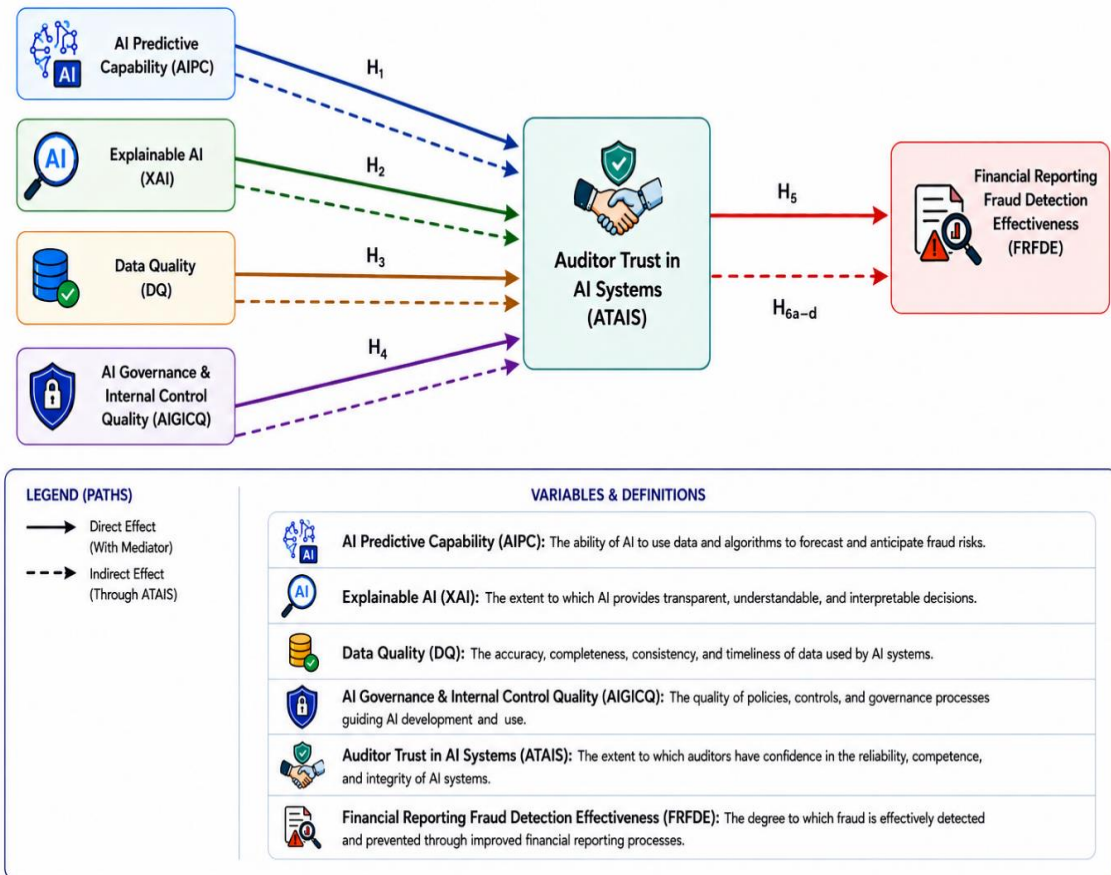


Figure 1: Conceptual framework of the study

2.9 Research Gaps

A careful analysis of this literature reveals several interrelated weaknesses that this study seeks to address. There is a practical shortcoming, as existing studies of fraud detection have predominantly focused on the performance of algorithms applied to archival data, while little research has investigated how qualified professionals evaluate, trust and monitor AI-enabled systems (Perols, 2011; Bao et al., 2020). There is a coherence gap, since AI capability, explainability, data quality, governance, trust and detection outcomes have evolved in largely independent research streams (Wang and Strong, 1996; Glikson and Woolley, 2020; Arrieta et al., 2020). There is a mechanism gap, because very little research investigates trust as a mediating mechanism linking these antecedents with detection outcomes (Glikson and Woolley, 2020; Raisch and Krakowski, 2021). There is a methodological gap, in that computational research produces technical proof with limited behavioural insight, while human–AI experiments yield evidence on trust with limited workforce relevance. Finally, there is a contextual gap, with very little integrated evidence examining these relationships among professionals in the U.S. financial reporting setting (Bao et al., 2020) — a setting of particular importance given the mounting interest in AI as a tool for addressing U.S. audit deficiencies and improving compliance with PCAOB and SEC standards (Yeboah et al., 2026) — whereas most existing integrated evidence on AI, audit quality and professional trust has been gathered in other institutional settings, such as emerging markets (Muhammad, 2025) and public sector reporting (Syahfir et al., 2024). This study addresses these deficiencies through the integrated, practice-based empirical model described above.

3. Research Methodology

3.1 Research Design and Analytical Approach

A quantitative, cross-sectional survey methodology was used to explore the factors influencing AI-augmented detection of financial reporting fraud in the U.S. financial reporting context. The proposed model specifies theoretically established relationships among latent constructs that can be measured with structured instruments and analysed using multivariate techniques, which made a quantitative methodology appropriate (Hair et al., 2022). Because cross-sectional data cannot establish temporal precedence, the identified relationships should be interpreted as theoretically grounded associations rather than definitive evidence of causation (Podsakoff et al., 2012).

The hypotheses were tested using Partial Least Squares Structural Equation Modelling (PLS-SEM) in SmartPLS version 4.1.1.6. PLS-SEM was appropriate for this complex model, which includes six latent constructs together with five direct and four indirect effects, given its emphasis on explaining variance in the endogenous constructs in a predictive mode (Hair et al., 2019, 2022). The assessment followed the standard two-stage procedure: first, the reflective measurement model was assessed through indicator loadings, internal consistency

reliability, convergent validity and discriminant validity; then the structural model was evaluated using collinearity diagnostics, R^2 , f^2 , path coefficients and bootstrapped significance tests for direct and specific indirect effects (Hair et al., 2019, 2022).

3.2 Measurement Instruments

Data were collected using a structured questionnaire of 24 items covering six latent constructs (four items per construct), adapted from previously validated instruments and tailored to the AI-assisted financial reporting fraud detection context (Table 1). All constructs were operationalised reflectively, and items were rated on a five-point Likert scale (1 = strongly disagree, 5 = strongly agree). The instrument was reviewed by experts in auditing, financial reporting and AI and pilot-tested with a small sample drawn from the target population before full-scale administration; following that evaluation, modifications were limited to wording and context specificity.

Table 1: Measurement instruments and their sources

Construct	Code	Items	Scale	Source
AI Predictive Capability	AIPC	4	5-point Likert	Bao et al. (2020); Cecchini et al. (2010)
Explainable AI	XAI	4	5-point Likert	Zhou et al. (2023)
Data Quality	DQ	4	5-point Likert	Wang and Strong (1996)
AI Governance and Internal Control Quality	AIGIC Q	4	5-point Likert	Tan et al. (2023)
Auditor Trust in AI Systems	ATAIS	4	5-point Likert	Glikson and Woolley (2020); Zhou et al. (2023)
Financial Reporting Fraud Detection Effectiveness	FRFDE	4	5-point Likert	Perols (2011); Bao et al. (2020)

3.3 Population and Sampling

The target population consisted of professionals in the U.S. financial reporting industry with relevant knowledge or experience in financial reporting, auditing, fraud detection, internal controls, compliance or AI-enhanced analytical technologies, including external and internal auditors, certified public accountants, forensic accountants, financial reporting specialists and compliance and risk management professionals. The unit of analysis was the individual professional. Screening criteria required participants to be currently based in the United States,

to hold relevant professional experience and to be familiar with AI-enhanced audit analytics or fraud detection technologies; respondents who did not meet these criteria were excluded.

Because no comprehensive sampling frame exists for U.S. professionals with AI experience relevant to the study domain, purposive sampling was used, sourcing participants from professional networks, accountancy and auditing groups, industry associations and relevant online professional platforms. Sample sufficiency was assessed following established guidelines for PLS-SEM sample-size determination (Cohen, 1988; Hair et al., 2022), taking into account that the maximum number of predictors directed at any endogenous construct is four (Auditor Trust in AI Systems) together with the anticipated effect size and required statistical power. The final sample comprised 450 valid responses.

Table 2: Demographic profile of respondents (N = 450)

Particulars	Category	Frequency	Percentage
Gender	Female	215	47.8
	Male	235	52.2
Age (years)	Below 20	9	2.0
	21-30	144	32.0
	31-40	253	56.2
	41 and above	44	9.8
Profession	Business	144	32.0
	Service	306	68.0
Income (USD)	Below 2,500	72	16.0
	2,500-3,499	135	30.0
	3,500-4,499	153	34.0
	4,500 and above	90	20.0

Note: N = 450. Frequencies and percentages were computed from the primary survey data file. Income is reported in U.S. dollars per month. Source: Survey data, 2026.

3.4 Data Collection and Procedure

Primary data were collected through a structured online survey, an approach appropriate to the national footprint of a target population distributed across the United States. The survey opened with an informed consent statement setting out the scholarly purpose of the study and confirming that participation was voluntary and confidential and that participants could withdraw at any time; no personally identifiable information was collected beyond what was required for research administration. Screening questions preceded the 24 substantive measurement items, and responses from ineligible persons were excluded from the final dataset.

The survey employed established procedural remedies to control for common method bias, including separating the predictor and criterion constructs within the instrument, carefully worded item formats, guaranteed respondent anonymity and an explicit reminder that there were no right or wrong answers (Podsakoff et al.,

2012). After data collection, the dataset was inspected for missing observations, incomplete responses and outlying response patterns before analysis. All measurement and structural model evaluation was performed using the two-step PLS-SEM procedure with bootstrapping (5,000 subsamples) to obtain standard errors, t-values and significance levels for the structural paths and the specific indirect (mediation) paths (Hair et al., 2019, 2022).

4. Results

4.1 Descriptive Statistics

Table 3 presents item-level descriptive statistics, reporting response tendencies, variability and distributional characteristics prior to model estimation (Hair et al., 2022). All 24 indicators had mean scores between 4.107 and 4.307, and all indicators had medians of either 4 or 5, indicating that participants generally expressed strong agreement with the indicators of AI Predictive Capability, Explainable AI, Data Quality, AI Governance and Internal Control Quality, Auditor Trust in AI Systems and Financial Reporting Fraud Detection Effectiveness. AIGICQ3 showed the highest mean ($M = 4.307$, $SD = 0.845$), followed closely by XAI3 ($M = 4.300$, $SD = 0.875$), while AIPC1 had the lowest ($M = 4.107$, $SD = 0.981$); even this minimum lies well above the scale midpoint. Standard deviations ranged from 0.832 to 0.981, suggesting moderate response variability with low dispersion. All items were negatively skewed (-1.142 to -0.505), indicating that responses were concentrated towards the upper end of the scale, and excess kurtosis ranged from -1.031 to 0.872 , indicating only minor deviations in peakedness. Cramér–von Mises tests returned $p < .001$ for every item. These represent small indicator-level deviations from univariate normality (Hair et al., 2019, 2022), but they are not a concern for the analysis, since PLS-SEM does not require multivariate normality and bootstrapping provides reliable significance tests of the path estimates.

Table 3: Descriptive statistics of the measurement items

Construct	Item	M	Mdn	SD	Excess kurtosis	Skewness	CvM p
AIPC	AIPC1	4.107	4.000	0.981	0.487	−0.953	<.001
	AIPC2	4.249	5.000	0.909	0.546	−1.064	<.001
	AIPC3	4.176	4.000	0.965	0.237	−0.983	<.001
	AIPC4	4.238	5.000	0.930	0.778	−1.121	<.001
XAI	XAI1	4.122	4.000	0.926	−0.393	−0.683	<.001
	XAI2	4.264	4.000	0.861	0.578	−1.020	<.001
	XAI3	4.300	5.000	0.875	0.872	−1.142	<.001
	XAI4	4.280	5.000	0.883	0.386	−1.026	<.001
DQ	DQ1	4.213	5.000	0.884	−1.021	−0.605	<.001
	DQ2	4.244	5.000	0.863	−0.831	−0.681	<.001
	DQ3	4.256	5.000	0.861	−0.643	−0.748	<.001
	DQ4	4.298	5.000	0.856	−0.724	−0.783	<.001
AIGICQ	AIGICQ1	4.269	4.000	0.836	−0.258	−0.837	<.001
	AIGICQ2	4.293	5.000	0.832	−0.408	−0.827	<.001
	AIGICQ3	4.307	5.000	0.845	0.069	−0.963	<.001
	AIGICQ4	4.291	5.000	0.863	−0.162	−0.933	<.001
ATAIS	ATAIS1	4.184	4.000	0.940	0.324	−0.971	<.001
	ATAIS2	4.256	5.000	0.921	0.307	−1.042	<.001
	ATAIS3	4.213	5.000	0.961	0.250	−1.011	<.001
	ATAIS4	4.291	5.000	0.896	−0.073	−0.979	<.001
FRFDE	FRFDE1	4.120	4.000	0.911	−1.031	−0.505	<.001
	FRFDE2	4.233	4.000	0.870	0.188	−0.918	<.001
	FRFDE3	4.256	5.000	0.871	−0.050	−0.905	<.001
	FRFDE4	4.260	5.000	0.875	−0.442	−0.831	<.001

Note: N = 450. AIPC = AI Predictive Capability; XAI = Explainable AI; DQ = Data Quality; AIGICQ = AI Governance and Internal Control Quality; ATAIS = Auditor Trust in AI Systems; FRFDE = Financial Reporting Fraud Detection Effectiveness; CvM = Cramér–von Mises. Source: SmartPLS output (version 4.1.1.6).

4.2 Common Method Bias

Full-collinearity variance inflation factors (VIFs), a key diagnostic in self-reported, cross-sectional survey research (Kock, 2015; Podsakoff et al., 2012), were used to check for common method bias and are reported in Table 4. The VIF values ranged from 1.000 to 2.727: 2.675 for AI Governance and Internal Control Quality, 2.460 for AI Predictive Capability, 1.000 for Auditor Trust in AI Systems, 2.727 for Data Quality and 2.412 for Explainable AI. All values remain well below the conservative full-collinearity threshold of 3.3 (Kock, 2015). While no single diagnostic can entirely rule out method bias (Podsakoff et al., 2012), the combination of established procedural remedies with full-collinearity values below the threshold indicates that pathological collinearity is unlikely in this dataset and that common method bias is not a material concern.

Table 4: Common method bias test (full-collinearity VIF)

Construct	Full-collinearity VIF
AI Governance and Internal Control Quality (AIGICQ)	2.675
AI Predictive Capability (AIPC)	2.460
Auditor Trust in AI Systems (ATAIS)	1.000
Data Quality (DQ)	2.727
Explainable AI (XAI)	2.412

Note: Threshold for full collinearity is 3.3 (Kock, 2015). Source: SmartPLS output (version 4.1.1.6).

4.3 Measurement Model Assessment

Table 5 reports the reliability and convergent validity of the reflective measurement model (Fornell and Larcker, 1981; Hair et al., 2022). All indicator loadings ranged between 0.871 and 0.922, indicating good indicator reliability, since all exceed the minimum threshold of 0.70. Cronbach's alpha ranged from 0.906 to 0.929 and composite reliability from 0.934 to 0.949, indicating sufficient internal consistency without redundancy in measurement. Average variance extracted (AVE) ranged from 0.781 to 0.824, well above the 0.50 threshold, confirming convergent validity for all six constructs (Fornell and Larcker, 1981). Item-level VIF values ranged from 2.399 to 3.755, acceptable under the commonly applied threshold of 5.0, although a few values exceed the stricter 3.0-3.3 rule and should be monitored in future refinements of the instrument (Hair et al., 2019, 2022).

Table 5: Construct validity and reliability

Construct	Item	Loading	Alpha	CR	AVE	VIF
AIGICQ	AIGICQ1	0.901	0.929	0.949	0.824	3.136
	AIGICQ2	0.922				3.755
	AIGICQ3	0.903				3.158
	AIGICQ4	0.905				3.233
AIPC	AIPC1	0.900	0.915	0.940	0.796	2.976
	AIPC2	0.893				2.911
	AIPC3	0.892				2.812
	AIPC4	0.884				2.751
ATAIS	ATAIS1	0.891	0.916	0.941	0.799	2.835
	ATAIS2	0.893				2.877
	ATAIS3	0.896				2.986
	ATAIS4	0.895				2.980
DQ	DQ1	0.901	0.922	0.945	0.811	3.054
	DQ2	0.902				3.158
	DQ3	0.890				2.837
	DQ4	0.909				3.354
FRFDE	FRFDE1	0.873	0.909	0.936	0.785	2.622
	FRFDE2	0.889				2.760
	FRFDE3	0.884				2.587
	FRFDE4	0.898				3.062
XAI	XAI1	0.885	0.906	0.934	0.781	2.720
	XAI2	0.889				2.790
	XAI3	0.871				2.399
	XAI4	0.889				2.817

Note: Thresholds are loading > .70, alpha > .70, CR > .70 and AVE > .50. CR = composite reliability; AVE = average variance extracted; VIF = variance inflation factor. Source: SmartPLS output (version 4.1.1.6).

Discriminant validity was assessed using the heterotrait–monotrait ratio (HTMT) and the Fornell–Larcker criterion (Fornell and Larcker, 1981; Henseler et al., 2015), reported in Tables 6 and 7. All HTMT ratios fell below the strict 0.85 threshold, with the highest at 0.831 between Explainable AI and Financial Reporting Fraud Detection Effectiveness and the lowest at 0.713 between AI Governance and Internal Control Quality and Financial Reporting Fraud Detection Effectiveness. Consistently, the square root of AVE for each construct (diagonal values ranging from 0.884 to 0.908) exceeded its correlations with all other constructs; for example, the value of 0.894 for Auditor Trust in AI Systems exceeds its highest inter-construct correlation of 0.743. Taken together, these findings indicate that the six constructs are empirically distinct while remaining interrelated.

Table 6: Discriminant validity based on the heterotrait–monotrait ratio (HTMT)

	AIGICQ	AIPC	ATAIS	DQ	FRFDE	XAI
AIGICQ						
AIPC	0.734					
ATAIS	0.806	0.792				
DQ	0.799	0.753	0.785			
FRFDE	0.713	0.791	0.791	0.740		
XAI	0.741	0.763	0.807	0.740	0.831	

Note: Values are reported below the diagonal. HTMT values below .85 support discriminant validity. Source: SmartPLS output (version 4.1.1.6).

Table 7: Discriminant validity based on the Fornell–Larcker criterion

	AIGICQ	AIPC	ATAIS	DQ	FRFDE	XAI
AIGICQ	0.908					
AIPC	0.677	0.892				
ATAIS	0.743	0.725	0.894			
DQ	0.740	0.691	0.721	0.900		
FRFDE	0.656	0.723	0.724	0.677	0.886	
XAI	0.680	0.694	0.736	0.677	0.753	0.884

Note: Diagonal elements are the square roots of the average variance extracted and exceed the off-diagonal correlations, supporting discriminant validity. Source: SmartPLS output (version 4.1.1.6).

4.4 Structural Model Assessment

The coefficients of determination for the endogenous constructs are reported in Table 8. The four antecedents jointly explained 69.7% of the variance in Auditor Trust in AI Systems ($R^2 = 0.697$; adjusted $R^2 = 0.694$), a strong level of explanatory power for a behavioural, organisational construct. Auditor trust alone explained 52.4% of the variance in Financial Reporting Fraud Detection Effectiveness ($R^2 = 0.524$; adjusted $R^2 = 0.523$). The negligible differences between R^2 and adjusted R^2 across constructs imply that explanatory power was not materially inflated by the number of predictors.

Table 8: Coefficients of determination (R^2)

Endogenous construct	R^2	R^2 adjusted
Auditor Trust in AI Systems (ATAIS)	0.697	0.694
Financial Reporting Fraud Detection Effectiveness (FRFDE)	0.524	0.523

Note: Source: SmartPLS output (version 4.1.1.6).

The effect sizes (f^2) in Table 9 indicate the contribution of each predictor when omitted from the model; following Cohen (1988) and Hair et al. (2022), values of 0.02, 0.15 and 0.35 denote small, medium and large effects respectively. The four antecedents — AI Governance and Internal Control Quality ($f^2 = 0.091$), AI Predictive Capability ($f^2 = 0.071$), Data Quality ($f^2 = 0.038$) and Explainable AI ($f^2 = 0.101$) — each exert a small effect on auditor trust, with the largest of the four attributable to Explainable AI. By contrast, the effect of auditor trust on fraud detection effectiveness was large ($f^2 = 1.100$), exceeding the conventional large-effect threshold. This implies that, although the antecedents contribute modestly on an individual basis, collectively they build trust, which in turn is the critical proximate driver of fraud detection effectiveness.

Table 9: Effect sizes (f^2)

Path	f^2
AIGICQ → ATAIS	0.091
AIPC → ATAIS	0.071
DQ → ATAIS	0.038
XAI → ATAIS	0.101
ATAIS → FRFDE	1.100

Note: Effect sizes of 0.02, 0.15 and 0.35 denote small, medium and large effects respectively (Cohen, 1988). Source: SmartPLS output (version 4.1.1.6).

Model-fit indicators (Table 10) serve as additional diagnostics rather than as a substitute for evaluating measurement quality, explanatory power and path estimation in PLS-SEM (Hair et al., 2019, 2022). SRMR values were 0.035 for the saturated model and 0.073 for the estimated model, both below the common cut-off of 0.08 and therefore acceptable as approximate fit. The normed fit index was 0.905 for the saturated model and 0.892 for the estimated model, indicating borderline incremental fit, with the estimated model falling marginally below the 0.90 benchmark. The discrepancy measures d_ULS and d_G and the chi-square statistic ($\chi^2 = 976.114$ for the saturated model and 1,112.697 for the estimated model) are sensitive to sample size, and bootstrap-derived HI95 and HI99 critical values for d_ULS and d_G were not computed, all of which warrant a cautious rather than definitive claim regarding model fit.

Table 10: Model fit indices

Fit index	Saturated model	Estimated model
SRMR	0.035	0.073
d _{ULS}	0.363	1.616
d _G	0.350	0.428
Chi-square	976.114	1,112.697
NFI	0.905	0.892

Note: SRMR = standardised root mean square residual; NFI = normed fit index. Source: SmartPLS output (version 4.1.1.6).

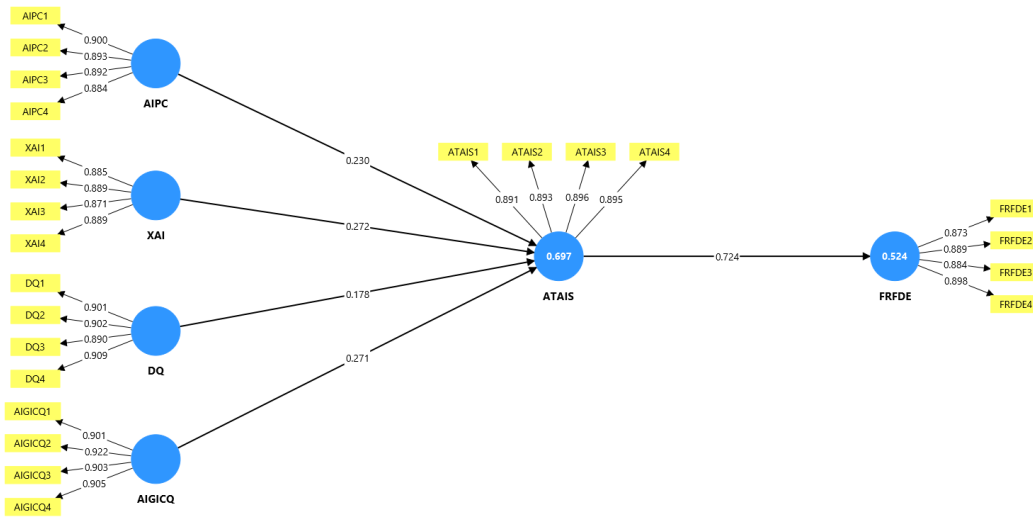
4.5 Hypothesis Testing

Table 11 reports the bootstrapped path coefficients and specific indirect effects (5,000 subsamples), ordered by hypothesis. AI predictive capability significantly increased auditor trust in AI systems ($\beta = 0.230$, $t = 3.489$, $p < .001$), supporting H1. Explainable AI exerted the strongest direct effect on auditor trust among the four antecedents ($\beta = 0.272$, $t = 3.839$, $p < .001$), supporting H2. Data quality had a smaller but significant positive effect ($\beta = 0.178$, $t = 2.795$, $p = .005$), supporting H3. The effect of AI governance and internal control quality on auditor trust was also significantly positive ($\beta = 0.271$, $t = 4.130$, $p < .001$), supporting H4. Auditor trust in AI systems, in turn, exerted a strong and significant direct effect on perceived financial reporting fraud detection effectiveness ($\beta = 0.724$, $t = 21.123$, $p < .001$), confirming H5 and supporting trust as the central mechanism of the model. All indirect effects through auditor trust in AI systems were positive and significant: AI predictive capability ($\beta = 0.166$, $t = 3.343$, $p = .001$), explainable AI ($\beta = 0.197$, $t = 3.690$, $p < .001$), data quality ($\beta = 0.129$, $t = 2.716$, $p = .007$) and AI governance and internal control quality ($\beta = 0.196$, $t = 4.337$, $p < .001$), supporting H6a to H6d. The largest indirect effect was produced by explainable AI, with AI governance and internal control quality a close second, while data quality produced the smallest but still significant indirect effect. Because the four antecedents were not specified with their own direct paths to fraud detection effectiveness, these results are described as significant mediation rather than as full or partial mediation. The estimated structural model is displayed in Figure 2.

Table 11: Results of hypothesis and mediation testing

Hyp.	Path	β	SD	t	p	Decision
H1	AIPC $\rightarrow$ ATAIS	0.230	0.066	3.489	<.001	Supported
H2	XAI $\rightarrow$ ATAIS	0.272	0.071	3.839	<.001	Supported
H3	DQ $\rightarrow$ ATAIS	0.178	0.064	2.795	0.005	Supported
H4	AIGICQ $\rightarrow$ ATAIS	0.271	0.066	4.130	<.001	Supported
H5	ATAIS $\rightarrow$ FRFDE	0.724	0.034	21.123	<.001	Supported
H6a	AIPC $\rightarrow$ ATAIS $\rightarrow$ FRFDE	0.166	0.050	3.343	0.001	Supported
H6b	XAI $\rightarrow$ ATAIS $\rightarrow$ FRFDE	0.197	0.053	3.690	<.001	Supported
H6c	DQ $\rightarrow$ ATAIS $\rightarrow$ FRFDE	0.129	0.047	2.716	0.007	Supported
H6d	AIGICQ $\rightarrow$ ATAIS $\rightarrow$ FRFDE	0.196	0.045	4.337	<.001	Supported

Note: Bootstrapping with 5,000 subsamples. β = standardised path coefficient; SD = standard deviation. Source: SmartPLS output (version 4.1.1.6).

**Figure 2: Structural model results (PLS algorithm, SmartPLS 4.1.1.6)**

5. Discussion

The findings imply that the contribution of advanced technology to the detection of financial reporting fraud is contingent on the interaction between algorithmic sophistication, data integrity, organisational governance and professional trust (Bao et al., 2020; Glikson and Woolley, 2020). The explanatory power of the model was substantial for both Auditor Trust in AI Systems (69.7%) and Perceived Financial Reporting Fraud Detection Effectiveness (52.4%), supporting the proposed framework (Hair et al., 2022) and reinforcing the view that AI creates organisational value only when the technology is combined with reliable data, established controls and informed professional judgment (Raisch and Krakowski, 2021).

Explainable AI was the most important predictor of trust among the four antecedents, indicating that explanations allow professionals to compare AI-generated suggestions with audit evidence, professional judgment and prior experience, although explanations must be designed carefully to avoid needless complexity or misleading simplicity (Arrieta et al., 2020; Glikson and Woolley, 2020). AI governance and internal control quality likewise had a relatively strong positive effect, indicating that auditors are more likely to trust AI when organisations maintain clear policies, accountability, oversight and monitoring (Shrestha et al., 2019); good governance reduces ambiguity about how AI systems are designed, managed and used in high-stakes decisions, suggesting that institutional safeguards still matter in AI-assisted decision-making (Raisch and Krakowski, 2021). Although AI predictive capability raised auditor trust considerably, its more modest effect supports the view that technical excellence alone cannot instil trust in the absence of transparency, quality data and governance (Arrieta et al., 2020), while remaining consistent with evidence that machine-learning algorithms identify complex fraud patterns more effectively than traditional methods (Cecchini et al., 2010; Bao et al., 2020). The direct effect of data quality was the smallest but nonetheless significant, confirming that accurate, complete, consistent and timely information is crucial for trustworthy AI-assisted analysis (Wang and Strong, 1996) and that professionals assess the credibility of AI partly on the basis of the integrity of the underlying data (Glikson and Woolley, 2020).

The strongest direct effect on fraud detection effectiveness in the model was that of auditor trust, confirming trust as the key behavioural driver (Glikson and Woolley, 2020): trust is necessary for professionals to assess, synthesise and apply machine-generated fraud indicators in their investigation and reporting decisions (Bao et al., 2020), supporting an automation–augmentation perspective in which the value of AI increases when it augments rather than substitutes for professional judgment and scepticism (Raisch and Krakowski, 2021). The mediation results reinforce this interpretation: auditor trust strongly mediated the effects of all four antecedents on fraud detection effectiveness, indicating that technological and organisational characteristics improve outcomes only when they are designed to support informed professional scepticism towards AI-generated evidence (Shrestha et al., 2019).

Collectively, the results show small direct antecedent effects on trust alongside a strong trust–effectiveness relationship, indicating that trust is not merely one antecedent among many but the medium through which the other conditions acquire practical significance.

6. Conclusion, Implications, Limitations and Future Research

6.1 Conclusion

The objective of this study was to examine the relative importance of technological, informational, organisational and behavioural factors in the use of AI for financial reporting fraud detection within the U.S. market. Auditor trust in AI systems was significantly predicted by AI predictive capability, explainable AI, data quality, and AI governance and internal control quality; it had a substantial positive effect on the effectiveness of financial reporting fraud detection and mediated the effects of all four antecedents. The implication is both plain and consequential: enhanced AI capability, better data, more interpretable explanations and stronger governance do not automatically improve fraud detection outcomes. They matter because they give professionals sufficient confidence to evaluate and deploy AI-generated fraud signals appropriately.

6.2 Theoretical Contribution

The study makes four theoretical contributions. First, by demonstrating that the professional efficacy of AI is contingent on explainability, data quality, governance, internal controls and auditor trust (Perols, 2011; Bao et al., 2020), it extends existing financial reporting fraud-detection research, which has largely concentrated on algorithm performance and classification results, and frames AI-enabled fraud detection as a socio-technical phenomenon (Raisch and Krakowski, 2021; Glikson and Woolley, 2020). Second, by establishing Auditor Trust in AI Systems as an explanatory node connecting organisational, technological and informational factors with perceptions of effectiveness in a contemporary, high-stakes professional domain that has received little empirical scrutiny, it develops human–AI trust theory (Glikson and Woolley, 2020; Lebovitz et al., 2022). Third, it extends the automation–augmentation perspective by showing that the organisational value of AI arises from the interaction between technological capability and professional judgment rather than from technological substitution, since AI-generated insights are still assessed and integrated by professionals (Raisch and Krakowski, 2021; Lebovitz et al., 2022). Finally, it integrates several literature streams into a single empirical model by conceptualising AI-enabled fraud detection as an enhancement of monitoring that supports professionals in addressing information asymmetry and identifying possible opportunistic reporting behaviour (Jensen and Meckling, 1976; Bao et al., 2020).

6.3 Managerial Implications

The results offer practical insights for financial institutions, corporate management, audit firms and other organisations deploying AI-based financial reporting fraud detection systems. Implementing AI should be regarded as a process of organisational transformation rather than as the acquisition of new software, because unless organisations simultaneously develop the informational, organisational and human conditions for successful use, investment in the latest technology is unlikely to deliver the anticipated fraud-detection benefits. Where system outputs will inform professional decisions, priority should be given to systems that combine strong predictive capability with understandable and traceable explanations, so that auditors can establish why a particular transaction or pattern has been flagged. Organisations should establish robust AI governance policies covering responsibility and accountability for AI-based systems, human oversight, ongoing system monitoring, access control, model validation and documentation, alongside continued investment in data governance covering ownership, validation, integration and monitoring. To complement the analytical strengths of AI with professional scepticism, contextual knowledge and human judgment, managers should invest in training, simulation and continuing professional education that enable audit teams to calibrate their trust and evaluate AI-generated evidence critically rather than accepting or rejecting it wholesale.

6.4 Limitations and Future Research

This study has certain limitations. By capturing perceptions at a single point in time, the cross-sectional design limits causal inference and prevents examination of how trust and perceived efficacy may change with continued use of the system. Although common method bias was assessed, the study relied on self-reported perception measures, which may be affected by response bias and subjective evaluation. Given the composition of the sample, the findings may not generalise to jurisdictions outside the United States or to settings with different regulatory contexts, levels of AI adoption and technological capability. Finally, because only four antecedents of auditor trust were considered, the model may not capture the full interplay of individual, technological, organisational and institutional factors relevant to trusted professional use of AI.

Future work should employ longitudinal designs to explore how trust, reliance and detection effectiveness evolve with experience. Survey data should also be paired with objective metrics such as actual detection accuracy, false-positive rates, audit outcomes and system usage statistics. Subsequent models may incorporate further antecedent, mediating or moderating variables, including AI literacy, professional scepticism, algorithm aversion, perceived AI risk, regulatory pressure and organisational AI readiness. Multi-industry and cross-country studies could establish whether these relationships generalise across regulatory, cultural and institutional contexts. Experimental, longitudinal and mixed-method studies would further advance understanding of how professionals construct calibrated trust and

of the most effective ways to combine human judgment with AI capability in the detection of financial reporting fraud.

DECLARATIONS

CRedit authorship contribution statement. Mohammad Musa Mia: writing – review and editing, writing – original draft, visualisation, validation, supervision, software, resources, project administration, methodology, investigation, formal analysis, data curation, conceptualisation. Md Zahurul Islam: writing – review and editing, writing – original draft, supervision, methodology, investigation, data curation, conceptualisation. Abu Sayed: writing – review and editing, writing – original draft, supervision, software, methodology, investigation, formal analysis, data curation, conceptualisation. Faysal Ahmed: writing – review and editing, writing – original draft, visualisation, validation, methodology, investigation, formal analysis, data curation, conceptualisation. Ataur Rahman: writing – review and editing, investigation, conceptualisation. Mohammad Ariful Aziz: investigation, formal analysis, data curation. Md Samirul Islam: writing – review and editing, investigation, conceptualisation.

Funding statement. No funding was received for this research.

Declaration of competing interest. The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Data availability. The data supporting the findings of this study are available from the corresponding author on reasonable request.

References

- [1] Ahmed, B. and Abdellah, L. (2025). L'intelligence artificielle au service de l'audit interne: Opportunités, défis et perspectives. Zenodo.
- [2] Antwi, B.O., Adelakun, B.O., Fatogun, D.T. and Olaiya, O.P. (2024). Enhancing audit accuracy: The role of AI in detecting financial anomalies and fraud. *Finance and Accounting Research Journal*, 6(6), pp. 1049 - 1068.
- [3] Arrieta, A.B., Díaz-Rodríguez, N., Del Ser, J., Bennetot, A., Tabik, S., Barbado, A., García, S., Gil-López, S., Molina, D., Benjamins, R., Chatila, R. and Herrera, F. (2020). Explainable artificial intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information Fusion*, 58, pp. 82 - 115.
- [4] Atf, Z. and Lewis, P.R. (2025). Is trust correlated with explainability in AI? A meta-analysis. *arXiv preprint*.
- [5] Bao, Y., Ke, B., Li, B., Yu, Y.J. and Zhang, J. (2020). Detecting accounting fraud in publicly traded U.S. firms using a machine learning approach. *Journal of Accounting Research*, 58(1), pp. 199 - 235.
- [6] Bauer, K., von Zahn, M. and Hinz, O. (2023). Expl(AI)ned: The impact of explainable artificial intelligence on users' information processing. *Information Systems Research*, 34(4), pp. 1582 - 1602.
- [7] Bondac, G., Stanescu, S., Ionescu, C.A., Duica, A. and Uzlău, M.C. (2026). Decision-making in complex systems using AI-based decision support: The role of trust, transparency, and data quality. *Electronics*, 15(2), Article 372.
- [8] Cecchini, M., Aytug, H., Koehler, G.J. and Pathak, P. (2010). Detecting management fraud in public companies. *Management Science*, 56(7), pp. 1146 - 1160.
- [9] Chen, J., Han, X. and Li, A.P. (2025). The algorithmic auditor: Machine learning in financial statement analysis and fraud detection. *Computer Science Bulletin*, 8(1), pp. 347 - 363.
- [10] Cohen, J. (1988). *Statistical power analysis for the behavioral sciences*. Second edition, Lawrence Erlbaum Associates, Hillsdale, NJ.
- [11] Dash, M., Princy, D., Kumar, M.S., Guntaj, J., Jain, R., Yadav, A. and Hashmi, S. (2025). Artificial intelligence in auditing: Enhancing fraud detection and risk assessment. *International Journal of Accounting and Economics Studies*, 12(SI-1), pp. 71 - 75.
- [12] Davis, F.D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), pp. 319 - 340.
- [13] De Araújo, D.C.J. (2025). Artificial intelligence and machine learning in fraud detection: Practical applications, predictive models, and ethical risks. *Revista Fisio & Terapia*, 29(150), pp. 22 - 23.
- [14] Dechow, P.M., Ge, W., Larson, C.R. and Sloan, R.G. (2011). Predicting material accounting misstatements. *Contemporary Accounting Research*, 28(1), pp. 17 - 82.

- [15] Duarte, B., Correia, F., Arriaga, P. and Paiva, A. (2022). AI trust: Can explainable AI enhance warranted trust? *Human Behavior and Emerging Technologies*, 2023, Article 4637678.
- [16] Fornell, C. and Larcker, D.F. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of Marketing Research*, 18(1), pp. 39 - 50.
- [17] Ganapathy, V. (2025). A comparative study of explainable artificial intelligence (XAI) techniques in financial auditing applications. *Edumania: An International Multidisciplinary Journal*, 3(3), pp. 185 - 215.
- [18] Gkegkas, M., Kydros, D. and Pazarskis, M. (2025). Using data analytics in financial statement fraud detection and prevention: A systematic review of methods, challenges, and future directions. *Journal of Risk and Financial Management*, 18(11), Article 598.
- [19] Glikson, E. and Woolley, A.W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of Management Annals*, 14(2), pp. 627 - 660.
- [20] Gu, H., Schreyer, M., Moffitt, K. and Vasarhelyi, M.A. (2024). Artificial intelligence co-piloted auditing. *International Journal of Accounting Information Systems*, 54, Article 100698.
- [21] Hair, J.F., Hult, G.T.M., Ringle, C.M. and Sarstedt, M. (2022). *A primer on partial least squares structural equation modeling (PLS-SEM)*. Third edition, Sage, Thousand Oaks, CA.
- [22] Hair, J.F., Risher, J.J., Sarstedt, M. and Ringle, C.M. (2019). When to use and how to report the results of PLS-SEM. *European Business Review*, 31(1), pp. 2 - 24.
- [23] Hastuti, D.A. and Widuri, R. (2025). Strengthening the effectiveness of accounting information systems: A systematic review of AI integration and internal control practices. *Edelweiss Applied Science and Technology*, 9(10), pp. 427 - 438.
- [24] Hatahali, S.K.A. (2025). Explainable AI (XAI) as a mechanism for trust and compliance. *Zenodo*.
- [25] Henseler, J., Ringle, C.M. and Sarstedt, M. (2015). A new criterion for assessing discriminant validity in variance-based structural equation modeling. *Journal of the Academy of Marketing Science*, 43(1), pp. 115 - 135.
- [26] Imane, L. (2025). Artificial intelligence in financial auditing: Improving efficiency and addressing ethical and regulatory challenges. *Brazilian Journal of Business*, 7(1), Article e76833.
- [27] Jensen, M.C. and Meckling, W.H. (1976). Theory of the firm: Managerial behavior, agency costs and ownership structure. *Journal of Financial Economics*, 3(4), pp. 305 - 360.
- [28] Khosravi, Y., Mazouei, N.R., Badiei, H. and Samadi, F. (2025). Identifying the factors influencing audit quality based on artificial intelligence and individual characteristics. *Journal of Empirical Research in Accounting*.

- [29] Kock, N. (2015). Common method bias in PLS-SEM: A full collinearity assessment approach. *International Journal of e-Collaboration*, 11(4), pp. 1 - 10.
- [30] Kokina, J., Blanchette, S., Davenport, T.H. and Pachamanova, D. (2025). Challenges and opportunities for artificial intelligence in auditing: Evidence from the field. *International Journal of Accounting Information Systems*, 56, Article 100734.
- [31] Law, K.K.F. and Shen, M. (2025). How does artificial intelligence shape audit firms? *Management Science*, 71(5), pp. 3641 - 3666.
- [32] Lebovitz, S., Levina, N. and Lifshitz-Assaf, H. (2022). To engage or not to engage with AI for critical judgments: How professionals deal with opacity when using AI for medical diagnosis. *Organization Science*, 33(1), pp. 126 - 148.
- [33] Mikalef, P. and Gupta, M. (2021). Artificial intelligence capability: Conceptualization, measurement calibration, and empirical study on its impact on organizational creativity and firm performance. *Information & Management*, 58(3), Article 103434.
- [34] Muhammad, H. (2025). The transformational impact of artificial intelligence on audit quality: An empirical study in an emerging market. *International Journal of Financial, Accounting, and Economics Studies*, 4(10), pp. 10 - 46.
- [35] Najm, B.T. (2025). From traditional audits to intelligent assurance: Exploring the value contribution of AI to financial control processes. Working paper.
- [36] Nugrahanti, T.P., Fardiman, F., Lanjarsih, L., Ratna, P. and Ritha, H. (2025). The role of explainable artificial intelligence in enhancing auditor judgment quality in Indonesia. *The ES Accounting and Finance*, 3(3), pp. 204 - 211.
- [37] Oyeniyi, L.D., Ugochukwu, C.E. and Mhlongo, N.Z. (2024). The influence of AI on financial reporting quality: A critical review and analysis. Zenodo.
- [38] Perols, J. (2011). Financial statement fraud detection: An analysis of statistical and machine learning algorithms. *Auditing: A Journal of Practice & Theory*, 30(2), pp. 19 - 50.
- [39] Podsakoff, P.M., MacKenzie, S.B. and Podsakoff, N.P. (2012). Sources of method bias in social science research and recommendations on how to control it. *Annual Review of Psychology*, 63, pp. 539 - 569.
- [40] Priyo, A. (2025). Transforming financial auditing in the era of artificial intelligence. *Global Management*, 2(2), pp. 33 - 37.
- [41] Raisch, S. and Krakowski, S. (2021). Artificial intelligence and management: The automation-augmentation paradox. *Academy of Management Review*, 46(1), pp. 192 - 210.
- [42] Raza, M., Qurashi, H., Haidar, A. and Raza, M.S. (2025). Artificial intelligence in auditing: Transforming fraud detection, risk assessment and assurance quality in financial reporting. *Journal of Asian Development Studies*, 14(3), pp. 453 - 466.

- [43] S, S., S, S. and Noorul, H.S. (2025). Transforming fraud detection in banking with explainable AI: Enhancing transparency and trust. *Journal of Technology Informatics and Engineering*, 4(2), pp. 251 - 260.
- [44] Sahla, W.A. and Latif, D.M. (2023). Perceptions of artificial intelligence (AI) usage on auditor judgment. *Indonesian Journal of Applied Accounting and Finance*, 3(2), pp. 91 - 104.
- [45] Shrestha, Y.R., Ben-Menahem, S.M. and von Krogh, G. (2019). Organizational decision-making structures in the age of artificial intelligence. *California Management Review*, 61(4), pp. 66 - 83.
- [46] Suyono, W.P., Puspa, E.S., Anugrah, S. and Firnanda, R. (2025). Redefining fraud detection: The synergy between auditor competency and AI-powered audit analytics. *RIGGS Journal of Artificial Intelligence and Digital Business*, 4(3), pp. 953 - 960.
- [47] Syahfir, H.A., Nirwana and Haliah (2024). Building public trust: The role of AI in preventing and exposing fraudulent financial reporting in the public sector: A systematic literature review. *West Science Accounting and Finance*, 2(3), pp. 519 - 535.
- [48] Tan, E., Jean, M.P., Simonofski, A., Tombal, T., Kleizen, B., Sabbe, M., Bechoux, L. and Willem, P. (2023). Artificial intelligence and algorithmic decisions in fraud detection: An interpretive structural model. *Data & Policy*, 5, Article e22.
- [49] Vardhan, N. (2021). Explainable AI for regulatory auditing and compliance in SAP financial systems. Preprint, Zenodo.
- [50] Vasarhelyi, M.A., Kogan, A. and Tuttle, B.M. (2015). Big data in accounting: An overview. *Accounting Horizons*, 29(2), pp. 381 - 396.
- [51] Wang, R.Y. and Strong, D.M. (1996). Beyond accuracy: What data quality means to data consumers. *Journal of Management Information Systems*, 12(4), pp. 5 - 33.
- [52] Yeboah, M.M., Agyei, E., Anim, B., Nwinyi, P. and Nartey, O.L. (2026). AI in financial auditing: Addressing U.S. audit failures and strengthening PCAOB/SEC compliance. Zenodo.
- [53] Zhou, Y., Li, H., Xiao, Z. and Qiu, J. (2023). A user-centered explainable artificial intelligence approach for financial fraud detection. *Finance Research Letters*, 58, Article 104309.

Appendix: Measurement Instrument

Table A1: Measurement items by construct

Construct	Item code	Item	Source
AI Predictive Capability	AIPC1	AI systems can identify unusual financial reporting patterns more accurately than traditional manual methods.	Bao et al. (2020)
	AIPC2	AI-based tools are effective in detecting anomalies that may indicate financial reporting fraud.	
	AIPC3	AI improves the ability of auditors or analysts to predict potential financial misstatements.	
	AIPC4	AI can process large volumes of financial data efficiently for fraud detection purposes.	
Explainable AI	XAI1	AI fraud detection systems should clearly explain how they reach their conclusions.	Glikson and Woolley (2020)
	XAI2	I am more likely to trust AI fraud detection results when the system provides understandable explanations.	
	XAI3	Explainable AI helps auditors or financial professionals evaluate suspicious financial reporting signals.	
	XAI4	AI-generated fraud alerts are more useful when users can trace the logic behind them.	
Data Quality	DQ1	Accurate financial data improves the effectiveness of AI-based fraud detection.	Wang and Strong (1996)
	DQ2	Complete accounting records are necessary for AI systems to detect financial reporting fraud.	
	DQ3	Consistent financial data improves the reliability of AI fraud detection outputs.	
	DQ4	Timely financial data helps AI systems identify fraud risks earlier.	
AI Governance and Internal Control Quality	AIGICQ1	Strong internal controls improve the effectiveness of AI-based fraud detection.	Tan et al. (2023)
	AIGICQ2	Clear organisational policies are necessary for responsible use of AI in financial reporting fraud detection.	
	AIGICQ3	Regular monitoring of AI systems improves the reliability of fraud detection outcomes.	
	AIGICQ4	Management support for AI governance strengthens fraud detection practices.	
Auditor Trust in AI Systems	ATAIS1	I trust AI systems to support financial reporting fraud detection.	Glikson and Woolley (2020)
	ATAIS2	I believe AI-generated fraud alerts are generally reliable.	
	ATAIS3	I feel confident using AI outputs when assessing financial reporting fraud risk.	

Construct	Item code	Item	Source
	ATAIS4	I would rely on AI recommendations when investigating suspicious financial reporting patterns.	
Financial Reporting Fraud Detection Effectiveness	FRFDE1	AI improves the accuracy of financial reporting fraud detection.	Bao et al. (2020)
	FRFDE2	AI reduces the time needed to identify potential financial reporting fraud.	
	FRFDE3	AI helps reduce human error in fraud detection procedures.	
	FRFDE4	AI strengthens the overall effectiveness of fraud detection in financial reporting.	

Note: All items were rated on a five-point Likert scale (1 = strongly disagree, 5 = strongly agree).