

DOI: <https://doi.org/10.38035/dijefa.v7i3><https://creativecommons.org/licenses/by/4.0/>

When Trust Sharpens Skepticism: Large Language Models in Indonesian Audit Practice

Kamal Amarullah^{1*}, Hamzah Ritchi², Ahmad Zakie Mubarrok³

¹Universitas Padjadjaran, Bandung, Indonesia, kamal23001@mail.unpad.ac.id

²Universitas Padjadjaran, Bandung, Indonesia, hamzah.ritchi@unpad.ac.id

³Universitas Padjadjaran, Bandung, Indonesia, ahmad.zakie@unpad.ac.id

*Corresponding Author: kamal23001@mail.unpad.ac.id¹

Abstract: This study examines the effects of perceived transparency, explainability, and social influence on auditors' trust in Large Language Models, and the effect of that trust on their professional skepticism at Indonesian public accounting firms. The gap between global acceptance and trust levels toward artificial intelligence systems suggests that technology adoption is not always matched by adequate evaluation, a condition relevant to auditors, who must remain critical toward Large Language Models given their tendency to produce inaccurate answers. This study used a quantitative approach with Partial Least Squares Structural Equation Modeling, involving 102 auditors selected through purposive sampling. Results show that all three antecedent variables positively and significantly affect auditors' trust in Large Language Models, with perceived transparency contributing most, while trust in Large Language Models also positively and significantly affects professional skepticism, a direction opposite to the reliance pattern reported in prior audit automation literature. These findings suggest that trust in artificial intelligence based technology can form in a calibrated manner, coexisting with auditors' awareness of system limitations rather than diminishing their professional skepticism.

Keywords: trust, large language model, professional skepticism, transparency, explainability social influence.

INTRODUCTION

Digital transformation has reshaped global business operating patterns over the past two decades, driving significant shifts in the social, economic, and cultural aspects of organizations (Surahman & Legowo, 2024). Trust functions as the foundation determining the sustainability of digital technology utilization within organizational processes, since without an adequate level of trust, technology adoption tends to remain at the technical usage stage rather than developing into substantive change in decision making (Capestro et al., 2024; Hooda et al., 2022; Kuen et al., 2023).

Artificial intelligence (AI) has emerged as one of the main pillars of business system modernization, defined as the capacity of computational systems to emulate human cognitive functions such as reasoning and decision making (Kahraman et al., 2024). The application of

AI improves operational efficiency and decision making quality through automation and service personalization (Akhunova et al., 2024). A 2025 survey by KPMG International and the University of Melbourne (Statista, 2025a), found that global trust in AI systems stood at 46 percent, well below its acceptance rate of 72 percent, indicating that technology adoption frequently occurs without adequate prior critical evaluation (Müller et al., 2025; Rana et al., 2024). This gap leads users to judge a technology as useful while remaining reluctant to rely on it fully, particularly in high risk tasks (Dhagarra et al., 2020; Kenesei et al., 2025).

This general pattern of adoption outpacing trust raises a sharper concern in professions where technology output directly informs judgment, such as financial statement auditing. The AI in audit market is projected to grow from approximately USD1.0 billion in 2023 to USD11.7 billion in 2033 (Market.us, 2024). The Big Four public accounting firms, namely Deloitte, PwC, EY, and KPMG, have integrated generative AI into audit practice, as exemplified by Deloitte's DARTbot, used by more than 18,000 audit professionals in the United States to accelerate audit risk identification (Deloitte, 2024a, 2024b). Large Language Models (LLMs), as a core component of such systems, are machine learning models trained on large volumes of text data to generate human-like text (Cerf, 2023; Zhao et al., 2023). Auditors at non-Big Four accounting firms have also made use of external LLMs, although the study by (Fotoh & Mugwira, 2025) shows that external LLMs are not yet able to produce audit reports that are comprehensive and contextually appropriate.

The characteristics inherent to LLMs across these practices also include a tendency to generate hallucinations, in the form of incorrect answers or answers unsupported by adequate evidence, even as these systems display impressive reasoning capabilities (Farquhar et al., 2024). Such hallucinations arise because LLM training mechanisms prioritize guessing over acknowledging the boundaries of the system's knowledge, such that LLMs are unable to distinguish correct information from information that merely appears convincing (Kalai et al., 2025). This condition renders auditor trust in LLMs an aspect with direct implications for audit quality, since auditors must assess the reliability of information sources used in the examination process (He et al., 2019). Auditors' level of trust in AI systems is believed to influence, in turn, their degree of reliance on that system's output (Commerford et al., 2024), and reliance not balanced by adequate understanding of the system's limitations and biases has the potential to weaken auditors' critical judgment of audit evidence (Murikah et al., 2024; Seethamraju & Hecimovic, 2023).

This uncritical reliance ultimately affects how auditors interpret audit evidence, which lies at the core of applying professional skepticism as required by auditing standards in the exercise of professional judgment (Hurt, 2010; IAASB, 2025). Lombardi et al. (2023) argue that the use of cognitive technologies such as LLMs without an adequate evaluative framework has the potential to worsen auditors' judgment bias, such that auditors require critical evaluation of technology output to avoid reliance on potentially biased information (Li, 2022). Auditor trust in LLMs is thus believed to be associated with the level of professional skepticism applied in assessing information and audit evidence.

Cognitive and affective processes shape trust in AI based systems (Chen et al., 2023). One such cognitive aspect is understanding of the system's characteristics, which leads to the concept of Explainable Artificial Intelligence (XAI) as an approach that renders AI's working process traceable and understandable to humans (Gunning & Aha, 2019). Transparency refers to users' capacity to understand how an AI system arrives at a given decision, while explainability refers to the system's capacity to provide a human understandable explanation grounded in the reasoning and logic underlying the decision produced (Arrieta et al., 2019; Shin & Park, 2019). Because evaluating both dimensions requires users to actively process and judge the system's underlying logic rather than respond to superficial cues, this study treats transparency and explainability as reflections of the central route to trust formation. Both

aspects have been empirically shown to be closely related to users' level of trust in AI systems, as understanding of the decision making mechanism provides a cognitive foundation for building trust grounded in rational understanding (Balasubramaniam et al., 2023; Shin, 2020a; Vössing et al., 2022).

Beyond cognitive considerations, trust in AI is also shaped by social dynamics and collective norms within the work environment, as support from colleagues and work groups can foster an individual's confidence in a given technology (Kim et al., 2024). Auditors working in environments with a culture of digital innovation tend to be more open to and trusting of technology use in supporting audit process quality (Leocádio et al., 2024; Perner, 2021; Vitali & Giuliani, 2024).

Indonesia, as the largest economy in Southeast Asia with 954 companies listed on the Indonesia Stock Exchange (Putri, 2024), constitutes a strategic empirical context for examining this mechanism, given that the value of Indonesia's GenAI technology market is projected to grow from USD355.97 million in 2025 to USD2.02 billion in 2030 (Statista, 2025b). Systematic data on the extent to which Indonesian auditors' trust in LLMs is formed, as well as the underlying psychological factors, remain very limited, underscoring the urgency of research examining this mechanism specifically within the context of Indonesian public accounting firms.

Prior research leaves gaps that warrant further investigation. (Fotoh & Mugwira, 2025) examined non Big Four auditors' perceptions of external LLM capacity without investigating the psychological factors underlying trust formation. Peters (2022) highlighted auditors' tendency to over rely on automated audit evidence, a finding that raises the question of whether a similar pattern occurs in LLM use. Kokina et al. (2025) identified transparency and explainability as key barriers to AI adoption without quantitatively testing their causal relationship with trust, while Nguyen et al. (2024) found auditor resistance to AI through an institutional theory approach without explaining the underlying psychological mechanism. These four studies indicate that the psychological mechanism through which auditor trust in LLMs is formed, via both cognitive and affective pathways, and its association with professional skepticism, has not yet been examined within a single integrated model.

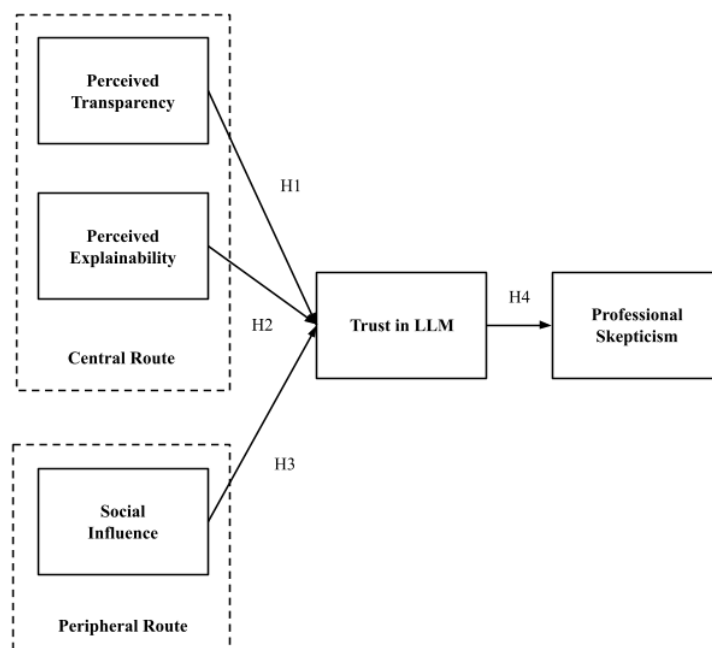


Figure 1. Research framework

This research contributes as a study applying the Elaboration Likelihood Model (ELM) within the domain of LLM based financial statement auditing, examining the effect of the central route, comprising perceived transparency and perceived explainability, and the peripheral route, comprising social influence, on auditor trust in LLMs and its effect on professional skepticism, among auditors at Indonesian public accounting firms using a Partial Least Squares Structural Equation Modeling (PLS-SEM) approach. Based on the foregoing, this research formulates the following hypotheses:

- H1: Perceived transparency, viewed as a reflection of the central route, has a positive effect on auditor trust in Large Language Models.
- H2: Perceived explainability, viewed as a reflection of the central route, has a positive effect on auditor trust in Large Language Models.
- H3: Social influence, viewed as a reflection of the peripheral route, has a positive effect on auditor trust in Large Language Models.
- H4: Auditor trust in Large Language Models has a significant effect on professional skepticism.

Literature Review

Elaboration Likelihood Model

The Elaboration Likelihood Model (ELM) is a persuasion theory developed by Petty and Cacioppo (1986) to explain how individuals process information and form attitudes through two main routes, namely the central route and the peripheral route. The central route is characterized by high cognitive involvement in critically evaluating arguments, producing attitudes that are stable and strongly predictive of behavior, whereas the peripheral route relies on surface level cues without in depth evaluation, producing attitudes that are more susceptible to change (Petty & Cacioppo, 1986; Tian et al., 2021; Zhou, 2012). This research positions perceived transparency and perceived explainability as reflections of the central route, since both require users to engage in argumentative evaluation based on system logic and mechanisms that can be critically examined (Arrieta et al., 2019; Ding et al., 2022; Felzmann et al., 2020). Social influence is positioned as a reflection of the peripheral route, since individuals tend to refer to the actions and views of others as a basis for decisions without direct evaluation of the argument itself (Chen et al., 2022; Rajak & Shaw, 2021; Zhou & Wu, 2025). The outcome of both persuasion routes is manifested in the form of attitude (Ding et al., 2025; Tian et al., 2021; Zhou, 2012), which in this research is operationalized as user trust in LLMs, such that ELM constitutes an appropriate framework for explaining how the three antecedent variables shape auditor trust through elaboration processes that differ in depth.

Attitude Theory

Attitude Theory, as described by Eagly and Chaiken (1993, 2007), views attitude as an evaluative tendency formed through cognitive, affective, and behavioral dimensions toward a given attitude object. This research employs Attitude Theory to explain how auditor trust in LLMs, as a form of attitude, affects the professional skepticism applied during the audit process. Attitude is understood as an evaluative tendency that helps individuals interpret their environment and maintain consistency of judgment, which subsequently shapes later responses (Cacioppo et al., 1989; Fazio, 1989; Greenwald, 1989). Auditor trust in LLMs is thus regarded as an evaluative tendency that influences how auditors process information generated by such systems, thereby affecting their degree of caution and vigilance in interpreting audit evidence, in line with auditing standards that affirm the role of professional skepticism in determining examination procedures during an engagement.

METHOD

Research Design and Sampling

This research employs a quantitative method with an explanatory approach to examine the effect of perceived transparency, perceived explainability, and social influence on auditor trust in Large Language Models, as well as the effect of that trust on professional skepticism. The study population comprises financial auditors working at public accounting firms in Indonesia, excluding junior auditors, who are considered not yet to possess sufficient experience to consistently apply professional skepticism. The sample was determined using purposive sampling, with criteria requiring auditors at the senior, supervisor, manager, or partner level who had used LLMs in financial statement audit engagements, with respondent distribution controlled so as not to be dominated by any single public accounting firm. The minimum sample size was determined using G*Power with an a priori approach, an effect size f^2 of 0.15, a significance level of 0.05, and a power of 0.80 for three predictors. A total of 102 respondents meeting these criteria were successfully collected through an online questionnaire, such that the obtained sample exceeded this threshold.

Measurement

The five constructs in this research are reflective and measured using a questionnaire with a 1 to 6 Likert scale. Perceived Transparency is measured through the dimensions of accuracy, clarity, and disclosure, adapted from Lee and Cha (2024) and Schnackenberg and Tomlinson (2016). Perceived Explainability is measured unidimensionally, adapted from Hamm et al. (2023) and Wang and Benbasat (2016). Social Influence is measured unidimensionally, adapted from Kim et al. (2024). Trust in LLM is measured through the dimensions of functionality, helpfulness, and reliability, adapted from Mcknight et al. (2011). Professional Skepticism is measured through the dimensions of questioning mind, suspension of judgment, and search for knowledge, adapted from items in Robinson et al. (2018), and constitutes a methodological contribution of this research, as these items were adapted specifically for the context of LLM use by financial auditors.

Data Analysis Methods

Data were analyzed using a Partial Least Squares Structural Equation Modeling (PLS-SEM) approach with SmartPLS 3 software. Evaluation of the measurement model (outer model) comprised convergent validity testing through outer loadings and Average Variance Extracted (AVE), discriminant validity testing through the Heterotrait Monotrait Ratio (HTMT), and reliability testing through Cronbach's Alpha and Composite Reliability (Hair et al., 2019, 2021). Evaluation of the structural model (inner model) comprised testing of the coefficient of determination (R^2), predictive relevance testing (Q^2) through the blindfolding procedure, effect size testing (F^2), and path significance testing through the bootstrapping procedure.

RESULTS AND DISCUSSION

Participants

This research analyzed 102 financial auditors working at public accounting firms (KAP) in Indonesia. Of this total, 55 respondents (54%) were male and 47 respondents (46%) were female, indicating a relatively balanced gender composition. Based on length of service, the majority of respondents, 85 individuals (83.34%), had between 2 and 6 years of experience, while the remainder had more than 6 years of service. Based on position, respondents were dominated by senior auditors, numbering 77 individuals (75.49%), followed by principals or managers, numbering 14 individuals (13.73%), supervisors, numbering 9 individuals (8.82%), and partners, numbering 2 individuals (1.96%).

The dominance of senior auditors in this sample indicates that the majority of respondents came from the executing level directly involved in technical audit work, including the use of LLMs during audit stages. The diversity of positions, ranging from senior auditor to partner, also reflects variation in the level of responsibility and experience in applying professional skepticism. Respondents were spread across 25 different public accounting firms, with 66.67 percent originating from firms affiliated with international networks, whether Big Four or non-Big Four, reflecting the purposive sampling strategy designed to avoid the dominance of any single firm.

Table 1. Profile of participants

Items	Values	Frequency	Percentage
Gender	Male	55	54%
	Female	47	46%
Work experience	2-6 years	85	83,34%
	> 6 years	17	16,66%
Role	Senior auditor	77	75,49%
	Supervisor	9	8,82%
	Principal/Manager	14	13,73%
	Partner	2	1,96%

Source: Research Data

Table 2. Purpose of LLM Use

Purpose of LLM Use	Frequency	% Respondents
As a consultation partner	69	67.65%
Assisting with audit planning	54	52.94%
Assisting with document analysis	53	51.96%
Assisting with analytical procedures	48	47.06%
Assisting with risk assessment	34	33.33%
Other	1	0.98%

Source: Research Data

Respondents were asked about their purposes for using LLMs in audit engagements, with the option to select more than one answer, such that the total number of responses collected reached 259 from 102 respondents. The most frequently selected purpose was as a consultation partner, chosen by 69 respondents (67.65%), followed by assistance with audit planning by 54 respondents (52.94%), assistance with document analysis by 53 respondents (51.96%), assistance with performing analytical procedures by 48 respondents (47.06%), and assistance with risk assessment by 34 respondents (33.33%). One respondent in the "other" category cited using LLMs as a means of obtaining industry or client company information during the pre engagement stage, as well as a means of searching for relevant accounting standards and regulatory references pertaining to issues encountered with the client under audit. The dominance of the consultation partner option indicates that auditors in this sample tend to position LLMs as a discussion partner in their thinking process, rather than merely as a single information retrieval tool.

Outer Model

Outer loading values reflect the strength of the relationship between an indicator and the construct it represents, with values above 0.708 recommended because they indicate that the construct is able to explain more than 50 percent of the variance in that indicator (Hair et al., 2021). Relatively weak loading values are commonly found in measurement instruments that

have only recently been developed, such that indicator deletion is not recommended to be performed automatically. Indicators with loading values between 0.40 and 0.708 are only considered for deletion if their removal is shown to meaningfully improve the construct's reliability or convergent validity, while still taking into account the implications for content validity, whereas indicators with loading values below 0.40 must be deleted without exception (Hair et al., 2021).

As presented in Table 3, outer loading values across the five constructs ranged from 0.542 to 0.878. Nine indicators, namely CLR2, DSC1, FNC2, RLB1, RLB4, QST1, QST3, SUS2, and SUS3, showed values below 0.708 but still above 0.40, and were therefore subjected to further evaluation of their impact on construct validity and reliability. The results of this evaluation showed that the AVE values, discriminant validity, and reliability of all constructs containing these indicators remained above the required thresholds, as presented in Table 4 through Table 6, such that all nine indicators were retained in the measurement model.

Table 3. Outer Loading

Variable	Indicators	Outer Loading
Perceived Transparency	ACC1	0,771
	ACC2	0,732
	ACC3	0,782
	ACC4	0,733
	CLR1	0,729
	CLR2	0,664
	DSC1	0,674
	DSC2	0,736
	DSC3	0,752
Perceived Explainability	EXP1	0,834
	EXP2	0,825
	EXP3	0,850
Social Influence	PS1	0,818
	PS2	0,789
	PS3	0,837
	PS4	0,849
	PS5	0,825
Trust in LLM	FNC1	0,817
	FNC2	0,683
	FNC3	0,731
	HLP1	0,745
	HLP2	0,816
	RLB1	0,647
	RLB2	0,778
	RLB3	0,727
Professional Skepticism	RLB4	0,636
	QST1	0,542
	QST2	0,758
	QST3	0,671
	SUS1	0,783
	SUS2	0,613
	SUS3	0,601
	SRC1	0,774
	SRC2	0,731
	SRC3	0,816
	SRC4	0,878

Source: Research Data

The AVE values for all constructs were above the minimum threshold of 0.50, ranging from 0.524 for Professional Skepticism to 0.700 for Perceived Explainability, as presented in Table 4.

Table 4. Average Variance Extracted (AVE)

Variable	Average Variance Extracted (AVE)
Perceived Transparency	0,543
Perceived Explainability	0,700
Social Influence	0,679
Trust in LLM	0,539
Professional Skepticism	0,524

Source: Research Data

Discriminant validity was tested using the Heterotrait Monotrait Ratio (HTMT) criterion, with a threshold of 0.90 (Hair et al., 2021). As presented in Table 5, all construct pairs showed HTMT values below this threshold, with the highest value of 0.894 between Trust in LLM and Perceived Transparency, and the lowest value of 0.276 between Professional Skepticism and Perceived Transparency. The low HTMT values between Professional Skepticism and the other constructs, ranging from 0.276 to 0.386, indicate that this construct is well differentiated from the three antecedent constructs of trust. Discriminant validity of the measurement model is therefore considered established, allowing the analysis to proceed to structural model testing.

Table 5. Heterotrait-Monotrait Ratio (HTMT)

Variable	Trust in LLM	Social Influence	Perceived Explainability	Perceived Transparency	Professional Skepticism
Trust in LLM					
Social Influence	0,802				
Perceived Explainability	0,861	0,759			
Perceived Transparency	0,894	0,734	0,808		
Professional Skepticism	0,361	0,315	0,386	0,276	

Source: Research Data

Table 6. Cronbach's Alpha and Composite Reliability

Variable	Cronbach's Alpha	Composite Reliability
Trust in LLM	0,892	0,913
Social Influence	0,882	0,913
Perceived Explainability	0,787	0,875
Perceived Transparency	0,894	0,914
Professional Skepticism	0,897	0,915

Source: Research Data

Reliability testing was conducted using Cronbach's Alpha and Composite Reliability values, with a threshold of 0.70 (Hair et al., 2021). As presented in Table 6, Cronbach's Alpha values ranged from 0.787 for Perceived Explainability to 0.897 for Professional Skepticism, while Composite Reliability values ranged from 0.875 to 0.915. All these values exceeded the minimum threshold, indicating that the research instrument possesses adequate internal consistency and can be relied upon to measure the constructs under study.

Inner Model

Figure 2 presents the results of the structural model (inner model) testing, displaying the path coefficients between the latent constructs in this research. Structural model testing comprised assessment of the R Square value to measure the extent to which the independent variables explain variation in the dependent variable, testing of path coefficients along with t statistic and p value to determine the significance of relationships between constructs, and calculation of F Square values to determine the relative contribution of each independent variable.

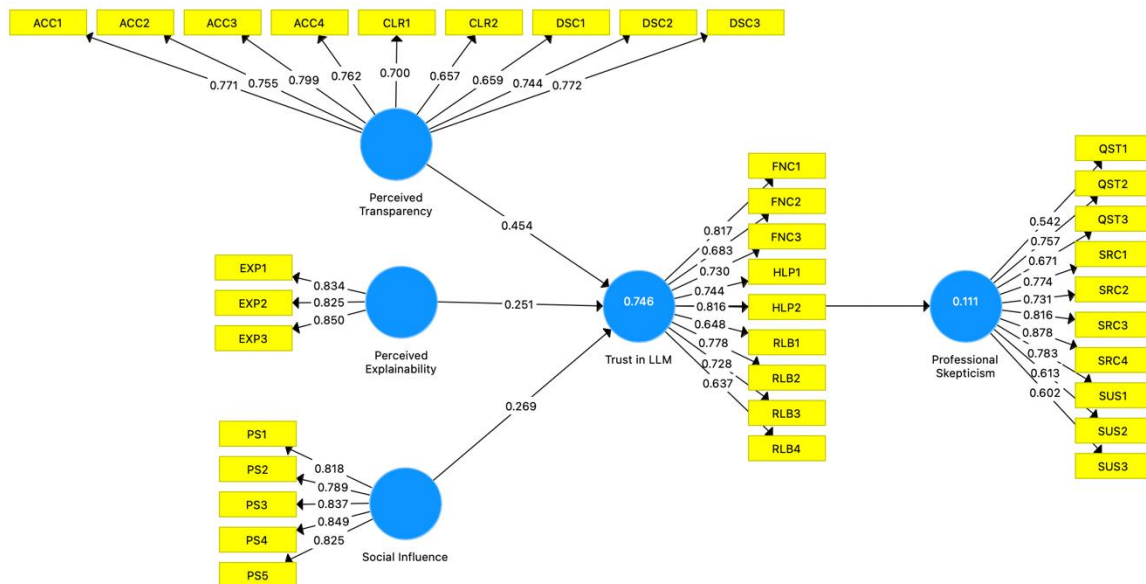


Figure 2. Inner Model

Structural model testing included the R Square (R^2) value to assess the capacity of the independent variables to explain variation in the dependent variable, with categories of 0.75, 0.50, and 0.25 respectively classified as strong, moderate, and weak (Hair et al., 2021). As presented in Table 7, trust in LLM had an adjusted R Square value of 0.739, approaching the strong category, while professional skepticism had an adjusted R Square value of 0.102, classified as weak. This contrast indicates that the three antecedent variables jointly explain a large proportion of the variation in auditor trust in LLM, while trust in LLM explains only a small portion of the variation in professional skepticism, given that professional skepticism is shaped by far more complex antecedent factors outside this model, including the professional orientation of the audit firm and social norms (Hardies et al., 2025).

Table 7. R-Square (R^2)

Variable	R-Square	R-Square Adjusted
Trust in LLM	0,746	0,739
Professional Skepticism	0,111	0,102

Source: Research Data

Predictive relevance (Q^2) values were obtained through the blindfolding procedure, with categories of 0.02, 0.15, and 0.35 respectively classified as small, medium, and large (Hair et al., 2021). Both dependent variables showed Q^2 values above 0, as presented in Table 8, with trust in LLM at 0.377 classified as large and professional skepticism at 0.046 classified as

small, such that this research model is considered to possess predictive relevance for both dependent variables.

Table 8. Predictive Relevance (Q²)

Variable	Trust in LLM	Social Influence	Perceived Explainability	Perceived Transparency	Professional Skepticism
Trust in LLM					0,125
Social Influence	0,143				
Perceived Explainability	0,115				
Perceived Transparency	0,359				
Professional Skepticism					

Source: Research Data

Effect size (F²) values were used to assess the individual contribution of each independent variable, with categories of 0.02, 0.15, and 0.35 respectively classified as small, medium, and large (Cohen, 1988). As presented in Table 9, perceived transparency showed the largest contribution to trust in LLM with an F² value of 0.359, followed by social influence at 0.143 and perceived explainability at 0.115, while trust in LLM showed an F² value of 0.125 for professional skepticism.

Table 9. Effect Size (F²)

Variable	SSO	SSE	Q ² (=1-SSE/SSO)
Trust in LLM	918,000	571,861	0,377
Professional Skepticism	1020,000	972,854	0,046

Source: Research Data

Hypothesis Testing

Hypothesis testing was conducted using path coefficients, t statistic values, and p values obtained from the bootstrapping procedure, with a relationship considered significant if the t statistic exceeded 1.96 and the p value was below 0.05. The results of testing the four hypotheses are presented in Table 10.

Table 10. Hypothesis Testing Results

Variable	Original Sample (O)	Sample Mean (M)	Standard Deviation (STDEV)	T Statistics (O/STDEV)	P Values
Trust in LLM -> Professional Skepticism	0,333	0,369	0,090	3,717	0,000
Social Influence -> Trust in LLM	0,269	0,277	0,081	3,335	0,001
Perceived Explainability -> Trust in LLM	0,251	0,249	0,084	2,992	0,003
Perceived Transparency -> Trust in LLM	0,454	0,449	0,093	4,889	0,000

Source: Research Data

All hypotheses in this research were accepted, with effects that were positive and statistically significant in direction. Perceived transparency showed the highest path coefficient (0.454) among the three antecedents of trust, indicating that auditors' perceptions of LLM transparency contribute most strongly to shaping their trust in the technology. Social influence showed a path coefficient of 0.269, while perceived explainability showed the lowest path

coefficient among the three antecedents, at 0.251. Trust in LLM was also shown to have a positive and significant effect on professional skepticism, with a path coefficient of 0.333, indicating that an increase in auditors' trust in LLM is associated with a corresponding increase in the professional skepticism exhibited in audit work.

Discussion

This discussion elaborates the results of testing the effect of perceived transparency, perceived explainability, and social influence on auditor trust in LLMs, as well as their implications for professional skepticism in audit practice. The gap between global acceptance and trust levels toward AI systems (Statista, 2025a) is relevant to examine within the audit profession, given that auditors are required to exercise careful professional judgment regarding the tools they use, including LLMs, which possess an inherent tendency toward hallucination. The three trust forming variables were selected based on the Elaboration Likelihood Model (ELM), which explains attitude formation through the central route and peripheral route, and Attitude Theory, which explains how trust, as a form of attitude, affects subsequent evaluative responses in the form of professional skepticism. The test results show that all independent variables have a positive and significant effect on auditor trust in LLM, and that trust also has a positive and significant effect on professional skepticism.

Perceived transparency showed a positive and significant effect on auditor trust in LLM, with the largest contribution among the three independent variables. This research views perceived transparency as a reflection of the central route within the ELM framework, since evaluating the accuracy, clarity, and sufficiency of information generated by LLMs requires active cognitive elaboration, rather than merely a passive response to surface level cues (Petty & Cacioppo, 1986). Trust formed through the central route is more stable and possesses stronger predictive power for subsequent behavior, such that auditor trust formed through perceived transparency is projected to be more robust than trust based on peripheral cues. This finding can be understood in light of LLMs' characteristics as black box systems whose internal mechanisms cannot be directly traced by users (Dam et al., 2018), such that perceived transparency becomes an important epistemic foundation for auditors in assessing the system's trustworthiness (von Eschenbach, 2021), explaining why its contribution is the largest among the three variables examined. This result is consistent with the findings of (Lee & Cha, 2024) on ChatGPT users and (Shin, 2020) in the context of recommendation algorithms, although both were studied on general user populations, as well as (Yu & Li, 2022) and (Sullivan & Weger, 2025), who examined human AI collaboration in work contexts, a setting functionally closer but not specific to the audit profession. This research complements that literature by showing that perceived transparency is a relevant determinant of trust among auditors as LLM users in financial statement audit practice.

Perceived explainability also showed a positive and significant effect on auditor trust in LLM, with a smaller contribution than perceived transparency. Similar to perceived transparency, perceived explainability is viewed as a reflection of the central route, since assessing the clarity of the explanation provided by an LLM for a given answer requires active reasoning to relate that explanation to the audit task context (Petty & Cacioppo, 1986). This finding can be understood in light of the role of understanding as an important prerequisite before trust can be meaningfully formed (Gefen et al., 2003; Gilpin et al., 2019), since system explanations enable auditors to interpret and evaluate LLM output, thereby reducing uncertainty regarding the system's reliability (Adadi & Berrada, 2018). This result is consistent with the findings of Hamm et al. (2023) in the context of signature forgery detection, Atf and Lewis (2025) through a cross domain meta analysis of AI applications, and Shin (2021), although the latter study positioned explainability as exerting its effect through transparency and accountability, in contrast to the present research, which positions it as a direct predictor.

This research complements that literature by demonstrating the relevance of perceived explainability as a determinant of trust in the context of auditors using LLMs in financial statement audit practice.

Meanwhile, social influence showed a positive and significant effect on auditor trust in LLM. Unlike the two preceding variables, this research views social influence as a reflection of the peripheral route, since auditors form judgments by referring to the views and actions of superiors and colleagues, without first conducting an in depth evaluation of LLM capability (Petty & Cacioppo, 1986). Trust formed through the peripheral route is more susceptible to change and possesses weaker predictive power than trust formed through the central route. This finding can be understood in light of individuals' tendency to interpret social support as a signal of a system's credibility (Mansoor, 2021), as well as their use of others' positive assessments as a reference point in forming beliefs (Oldeweme et al., 2021; Zhang et al., 2020). Auditors working within public accounting firms tend to treat organizational views and guidance from senior colleagues as reference points when using LLMs, such that trust can be formed through the social pathway without in depth technical evaluation. This result is consistent with the findings of Hooda et al. (2022) in the e government context, Rathnayake et al. (2025) on the adoption of government chatbots, and Heo and Na (2025) on LLM adoption in the construction industry, although the latter study examined the effect of social influence on intention to use rather than on trust directly. This research complements that literature by demonstrating the relevance of social influence as a determinant of trust in the context of auditors using LLMs.

Furthermore, trust in LLM showed a positive and significant effect on auditors' professional skepticism. Attitude Theory is used to explain this mechanism, in which attitude is understood as an evaluative tendency that influences subsequent responses (Cacioppo et al., 1989; Fazio, 1989; Greenwald, 1989), such that auditor trust in LLM shapes the way they process information generated by the system, which in turn affects the professional skepticism applied in interpreting audit evidence, consistent with auditing standards that affirm the importance of professional skepticism throughout an engagement (IAASB, 2025).

This pattern is noteworthy to examine, given that trust in technology is generally associated with reliance that has the potential to weaken critical evaluation (He et al., 2023; Klingbeil et al., 2024; Küper & Krämer, 2024). In the audit context, this kind of reliance pattern was previously identified by Peters (2022), who found a tendency among auditors to over rely on audit evidence from automation tools, prompting Fotoh and Mugwira (2025) to question whether a similar pattern emerges when auditors use LLMs. The results of this research show the opposite direction, auditor trust in LLM is instead positively associated with professional skepticism, rather than negatively associated as in the reliance pattern found by Peters (2022) regarding other automation tools.

This characteristic is reflected in how auditors make use of LLMs, where the majority of respondents used LLMs as a consultation partner, treating them as a discussion partner in their thinking process rather than as a producer of final conclusions accepted outright. The hallucination tendency inherent to LLMs, namely the system's tendency to guess rather than acknowledge a lack of knowledge, thereby producing answers that appear convincing yet are incorrect (Farquhar et al., 2024; Kalai et al., 2025), also constitutes a possible explanation, as awareness of this limitation may allow auditor trust to form alongside an understanding of its boundaries. Ding et al. (2026) found that the effect of reliance on vigilance depends on how trust is calibrated, such that reliance arising from calibrated trust does not always reduce user vigilance, opening the possibility that a similar pattern applies in the context of financial statement auditing.

To the extent of the literature search conducted, no prior research has directly examined the relationship between trust in LLM and auditors' professional skepticism, such that this finding provides an initial contribution to an area still rarely examined empirically, given that

the audit technology literature has more often focused on usage intention than on professional skepticism responses (Afsay et al., 2023; Anh et al., 2024; Ferri et al., 2020).

CONCLUSION

This research aims to examine the effect of perceived transparency, perceived explainability, and social influence on auditor trust in Large Language Models (LLM), as well as the effect of that trust on the professional skepticism of auditors at Indonesian public accounting firms. The test results show that all three antecedents, whether operating through the central route in the form of perceived transparency and explainability, or the peripheral route in the form of social influence, consistently exert a positive and significant effect on auditor trust in LLM, with perceived transparency providing the largest contribution. Trust in LLM was subsequently shown to have a positive and significant effect on professional skepticism, indicating that auditor trust in this technology does not reduce their vigilance in interpreting audit evidence, but instead moves in tandem with increased professional skepticism, although its contribution to the variability of skepticism is small, such that other factors beyond trust also play an important role.

This finding contributes to the audit literature by presenting initial empirical evidence applying the Elaboration Likelihood Model to explain the formation of auditor trust in LLM as well as its association with professional skepticism, a relationship whose direction differs from the reliance pattern reported in the audit literature on other automation tools. This result indicates that trust in AI based technology can be formed in a calibrated manner, coexisting with auditors' awareness of the system's limitations, rather than replacing their professional judgment.

Interpretation of this finding should take into account the cross sectional design employed, which limits conclusions regarding how auditor trust and skepticism develop over time, as well as the use of a purposive sample of 102 auditors, which limits the generalization of results to the broader auditor population. Trust in LLM was also measured as a single construct without distinguishing its cognitive and affective dimensions, while the classification of the central route and peripheral route in this research represents a theoretical position that has not been directly tested through separate measurement of auditors' motivation and cognitive elaboration.

Public accounting firms are advised to strengthen perceived transparency and explainability as criteria in using LLMs within the audit process, while also actively encouraging leadership support to accelerate the formation of audit teams' trust in this technology, without reducing ongoing professional skepticism training. Future research is advised to employ a longitudinal or experimental design to more adequately test the trust calibration process, to measure the cognitive and affective dimensions of trust separately, and to explore other antecedent factors of professional skepticism beyond trust in technology.

REFERENCES

- Adadi, A., & Berrada, M. (2018). Peeking Inside the Black-Box: A Survey on Explainable Artificial Intelligence (XAI). *IEEE Access*, 6, 52138–52160. <https://doi.org/10.1109/ACCESS.2018.2870052>
- Afsay, A., Tahriri, A., & Rezaee, Z. (2023). A meta-analysis of factors affecting acceptance of information technology in auditing. *International Journal of Accounting Information Systems*, 49, 100608. <https://doi.org/10.1016/j.accinf.2022.100608>
- Akhunova, S., Tuychiyeva, O., Tukhtasinova, D., & Akhunova, M. (2024). An innovative navigating system design along with AI tech integration for modern system implementation for business. *2024 4th International Conference on Advance Computing*

- and Innovative Technologies in Engineering (ICACITE), 462–467. <https://doi.org/10.1109/icacite60783.2024.10616591>
- Anh, N. T. M., Hoa, L. T. K., Thao, L. P., Nhi, D. A., Long, N. T., Truc, N. T., & Ngoc Xuan, V. (2024). The Effect of Technology Readiness on Adopting Artificial Intelligence in Accounting and Auditing in Vietnam. *Journal of Risk and Financial Management*, 17(1), 27. <https://doi.org/10.3390/jrfm17010027>
- Arrieta, A. B., Díaz-Rodríguez, N., Del Ser, J., Bennetot, A., Tabik, S., Barbado, A., Garcia, S., Gil-Lopez, S., Molina, D., Benjamins, R., Chatila, R., & Herrera, F. (2019). Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information Fusion*, 58, 82–115. <https://doi.org/10.1016/j.inffus.2019.12.012>
- Atf, Z., & Lewis, P. R. (2025). Is Trust Correlated With Explainability in AI? A Meta-Analysis. *IEEE Transactions on Technology and Society*, 7(1), 70–77. <https://doi.org/10.1109/TTS.2025.3558448>
- Balasubramaniam, N., Kauppinen, M., Rannisto, A., Hiekkanen, K., & Kujala, S. (2023). Transparency and explainability of AI systems: From ethical guidelines to requirements. *Information and Software Technology*, 159, 107197. <https://doi.org/10.1016/j.infsof.2023.107197>
- Cacioppo, J. T., Petty, R. E., & Geen, T. R. (1989). Attitude Structure and Function: From the Tripartite to the Homeostasis Model of Attitudes. In *Attitude Structure and Function*. Psychology Press.
- Capestro, M., Rizzo, C., Kliestik, T., Peluso, A. M., & Pino, G. (2024). Enabling digital technologies adoption in industrial districts: The key role of trust and knowledge sharing. *Technological Forecasting and Social Change*, 198, 123003. <https://doi.org/10.1016/j.techfore.2023.123003>
- Cerf, V. G. (2023). Large language models. *Communications of the ACM*, 66(8), 7–7. <https://doi.org/10.1145/3606337>
- Chen, C.-J., Tsai, P.-H., & Tang, J.-W. (2022). How informational-based readiness and social influence affect usage intentions of self-service stores through different routes: An elaboration likelihood model perspective. *Asia Pacific Business Review*, 28(3), 380–409. <https://doi.org/10.1080/13602381.2021.1872912>
- Chen, Q., Yin, C., & Gong, Y. (2023). Would an AI chatbot persuade you: An empirical answer from the elaboration likelihood model. *Information Technology & People*, 38(2), 937–962. <https://doi.org/10.1108/ITP-10-2021-0764>
- Cohen, J. (1988). *Statistical power analysis for the behavioral sciences* (2nd ed). L. Erlbaum Associates.
- Commerford, B., Eilifsen, A., Hatfield, R. C., Holmstrom, K. M., & Kinserdal, F. (2024). *Control Issues: How Providing Input Affects Auditors' Reliance on Artificial Intelligence* (SSRN Scholarly Paper No. 4847775). Social Science Research Network. <https://doi.org/10.1111/1911-3846.12974>
- Dam, H. K., Tran, T., & Ghose, A. (2018). Explainable software analytics. *Proceedings of the 40th International Conference on Software Engineering: New Ideas and Emerging Results, ICSE-NIER '18*, 53–56. <https://doi.org/10.1145/3183399.3183424>
- Deloitte. (2024a). Generative AI in accounting: Opportunities and risks to assess today. Deloitte. <https://www.deloitte.com/us/en/services/audit/blogs/accounting-finance/generative-ai-auditing.html>
- Deloitte. (2024b, November 13). Questions about implementing GenAI? Deloitte provides insights from its own AI journey. <https://www2.deloitte.com/us/en/blog/accounting-finance-blog/2024/answering-ai-questions-in-audit-assurance.html>

- Dhagarra, D., Goswami, M., & Kumar, G. (2020). Impact of Trust and Privacy Concerns on Technology Acceptance in Healthcare: An Indian Perspective. *International Journal of Medical Informatics*, 141, 104164. <https://doi.org/10.1016/j.ijmedinf.2020.104164>
- Ding, R., Chen, X., Wei, S., & Wang, J. (2025). What drives trust building in live streaming E-commerce? From an elaboration likelihood model perspective. *Industrial Management & Data Systems*, 125(3), 969–999. <https://doi.org/10.1108/IMDS-03-2024-0273>
- Ding, W., Abdel-Basset, M., Hawash, H., & Ali, A. M. (2022). Explainability of artificial intelligence methods, applications and challenges: A comprehensive survey. *Information Sciences*, 615, 238–292. <https://doi.org/10.1016/j.ins.2022.10.013>
- Ding, W., Wu, Y., & Wang, J. (2026). Beyond Risk Reduction: Vigilant Trust in Artificial Intelligence Based on Evidence from China. *Behavioral Sciences*, 16(1), 95. <https://doi.org/10.3390/bs16010095>
- Eagly, A. H., & Chaiken, S. (1993). *The psychology of attitudes* (pp. xxii, 794). Harcourt Brace Jovanovich College Publishers.
- Eagly, A. H., & Chaiken, S. (2007). The Advantages of an Inclusive Definition of Attitude. *Social Cognition*, 25(5), 582–602. <https://doi.org/10.1521/soco.2007.25.5.582>
- Farquhar, S., Kossen, J., Kuhn, L., & Gal, Y. (2024). Detecting hallucinations in large language models using semantic entropy. *Nature*, 630(8017), 625–630. <https://doi.org/10.1038/s41586-024-07421-0>
- Fazio, R. H. (1989). On the Power and Functionality of Attitudes: The Role of Attitude Accessibility. In *Attitude Structure and Function*. Psychology Press.
- Felzmann, H., Fosch-Villaronga, E., Lutz, C., & Tamò-Larrieux, A. (2020). Towards Transparency by Design for Artificial Intelligence. *Science and Engineering Ethics*, 26(6), 3333–3361. <https://doi.org/10.1007/s11948-020-00276-4>
- Ferri, L., Spanò, R., Ginesti, G., & Theodosopoulos, G. (2020). Ascertaining auditors' intentions to use blockchain technology: Evidence from the Big 4 accountancy firms in Italy. *Meditari Accountancy Research*, 29(5), 1063–1087. <https://doi.org/10.1108/MEDAR-03-2020-0829>
- Fotoh, E., & Mugwira, T. (2025). *Exploring Large Language Models in External Audits: Implications and Ethical Considerations*. SSRN. <https://doi.org/10.2139/ssrn.4932003>
- Gefen, Karahanna, & Straub. (2003). Trust and TAM in Online Shopping: An Integrated Model. *MIS Quarterly*, 27(1), 51. <https://doi.org/10.2307/30036519>
- Gilpin, L. H., Testart, C., Fruchter, N., & Adebayo, J. (2019). *Explaining Explanations to Society* (arXiv:1901.06560). arXiv. <https://doi.org/10.48550/arXiv.1901.06560>
- Greenwald, A. G. (1989). Why Attitudes are Important: Defining Attitude and Attitude Theory 20 Years Later. In *Attitude Structure and Function*. Psychology Press.
- Gunning, D., & Aha, D. W. (2019). DARPA's Explainable Artificial Intelligence Program. *AI Magazine*, 40(2), 44–58. <https://doi.org/10.1609/aimag.v40i2.2850>
- Hair, J. F., Hult, G. T. M., Ringle, C. M., Sarstedt, M., Danks, N. P., & Ray, S. (2021). *Partial Least Squares Structural Equation Modeling (PLS-SEM) Using R: A Workbook*. Springer International Publishing. <https://doi.org/10.1007/978-3-030-80519-7>
- Hamm, P., Klesel, M., Coberger, P., & Wittmann, H. F. (2023). Explanation matters: An experimental study on explainable AI. *Electronic Markets*, 33(1), 17. <https://doi.org/10.1007/s12525-023-00640-9>
- Hardies, K., Janssen, S., Vanstraelen, A., & Zehms, K. M. (2025). Using Field-Based Evidence to Understand the Antecedents to Auditors' Skeptical Actions. *AUDITING: A Journal of Practice & Theory*, 44(1), 105–135. <https://doi.org/10.2308/AJPT-2021-141>
- He, G., Kuiper, L., & Gadiraju, U. (2023). Knowing About Knowing: An Illusion of Human Competence Can Hinder Appropriate Reliance on AI Systems. *Proceedings of the 2023*

- CHI Conference on Human Factors in Computing Systems*, 1–18. <https://doi.org/10.1145/3544548.3581025>
- He, W., Sidhu, B., & Taylor, S. (2019). Audit quality and properties of analysts' information environment. *Journal of Business Finance & Accounting*, 46(3–4), 400–419. <https://doi.org/10.1111/jbfa.12358>
- Heo, S., & Na, S. (2025). Ready for departure: Factors to adopt large language model (LLM)-based artificial intelligence (AI) technology in the architecture, engineering and construction (AEC) industry. *Results in Engineering*, 25, 104325. <https://doi.org/10.1016/j.rineng.2025.104325>
- Hooda, A., Gupta, P., Jeyaraj, A., Giannakis, M., & Dwivedi, Y. K. (2022). The effects of trust on behavioral intention and use behavior within e-government contexts. *International Journal of Information Management*, 67, 102553. <https://doi.org/10.1016/j.ijinfomgt.2022.102553>
- Hurtt, R. K. (2010). Development of a Scale to Measure Professional Skepticism. *AUDITING: A Journal of Practice & Theory*, 29(1), 149–171. <https://doi.org/10.2308/aud.2010.29.1.149>
- IAASB. (2025). *Handbook of International Quality Management, Auditing, Review, Other Assurance, and Related Services Pronouncements*. www.iaasb.org
- Kahraman, F., Aktas, A., Bayrakceken, S., Çakar, T., Tarcan, H. S., Bayram, B., Durak, B., & Ulman, Y. I. (2024). Physicians' ethical concerns about artificial intelligence in medicine: A qualitative study: "The final decision should rest with a human." *Frontiers in Public Health*, 12. <https://doi.org/10.3389/fpubh.2024.1428396>
- Kalai, A. T., Nachum, O., Vempala, S. S., & Zhang, E. (2025). *Why Language Models Hallucinate* (arXiv:2509.04664). arXiv. <https://doi.org/10.48550/arXiv.2509.04664>
- Kenesei, Z., Kökény, L., Ásványi, K., & Jászberényi, M. (2025). The central role of trust and perceived risk in the acceptance of autonomous vehicles in an integrated UTAUT model. *European Transport Research Review*, 17(1), 8. <https://doi.org/10.1186/s12544-024-00681-x>
- Kim, Y., Blazquez, V., & Oh, T. (2024). Determinants of Generative AI System Adoption and Usage Behavior in Korean Companies: Applying the UTAUT Model. *Behavioral Sciences*, 14(11), 1035. <https://doi.org/10.3390/bs14111035>
- Kim, Y. J., Choi, J. H., & Fotso, G. M. N. (2024). Medical professionals' adoption of AI-based medical devices: UTAUT model with trust mediation. *Journal of Open Innovation: Technology, Market, and Complexity*, 10(1), 100220. <https://doi.org/10.1016/j.joitmc.2024.100220>
- Klingbeil, A., Grützner, C., & Schreck, P. (2024). Trust and reliance on AI — An experimental study on the extent and costs of overreliance on AI. *Computers in Human Behavior*, 160, 108352. <https://doi.org/10.1016/j.chb.2024.108352>
- Kokina, J., Blanchette, S., Davenport, T. H., & Pachamanova, D. (2025). Challenges and opportunities for artificial intelligence in auditing: Evidence from the field. *International Journal of Accounting Information Systems*, 56, 100734. <https://doi.org/10.1016/j.accinf.2025.100734>
- Kuen, L., Westmattmann, D., Bruckes, M., & Schewe, G. (2023). Who earns trust in online environments? A meta-analysis of trust in technology and trust in provider for technology acceptance. *Electronic Markets*, 33(1), 61. <https://doi.org/10.1007/s12525-023-00672-1>
- Küper, A., & Krämer, N. (2024). Psychological Traits and Appropriate Reliance: Factors Shaping Trust in AI. *International Journal of Human-Computer Interaction*, 1–17. <https://doi.org/10.1080/10447318.2024.2348216>

- Lee, C., & Cha, K. (2024). Toward the Dynamic Relationship Between AI Transparency and Trust in AI: A Case Study on ChatGPT. *International Journal of Human–Computer Interaction*, 41(13), 8086–8103. <https://doi.org/10.1080/10447318.2024.2405266>
- Leocádio, D., Malheiro, L., & Reis, J. C. G. dos. (2024). Auditors in the digital age: A systematic literature review. *Digital Transformation and Society*, 4(1), 5–20. <https://doi.org/10.1108/DTS-02-2024-0014>
- Li, X. (2022). Behavioral challenges to professional skepticism in auditors' data analytics journey. *Maandblad Voor Accountancy En Bedrijfseconomie*, 96(1/2), 45–54. <https://doi.org/10.5117/mab.96.78525>
- Lombardi, D. R., Sipior, J. C., & Dannemiller, S. (2023). Auditor Judgment Bias Research: A 50-Year Trend Analysis and Emerging Technology Use. *Journal of Information Systems*, 37(1), 109–141. <https://doi.org/10.2308/ISYS-2020-079>
- Mansoor, M. (2021). Citizens' trust in government as a function of good governance and government agency's provision of quality information on social media during COVID-19. *Government Information Quarterly*, 38(4), 101597. <https://doi.org/10.1016/j.giq.2021.101597>
- Market.us. (2024). *Global AI in Audit Market*. <https://market.us/report/ai-in-audit-market/>
- Mcknight, D. H., Carter, M., Thatcher, J. B., & Clay, P. F. (2011). Trust in a specific technology: An investigation of its components and measures. *ACM Transactions on Management Information Systems*, 2(2), 1–25. <https://doi.org/10.1145/1985347.1985353>
- Müller, A. P. R., Tangi, L., & Lerusse, A. (2025). Understanding the Adoption of Artificial Intelligence in Local Government Decision-Making: The Influence of Institutional Pressures and Managerial Perceptions. *Public Administration*, n/a(n/a). <https://doi.org/10.1111/padm.70033>
- Murikah, W., Nthenge, J. K., & Musyoka, F. M. (2024). Bias and ethics of AI systems applied in auditing—A systematic review. *Scientific African*, 25, e02281. <https://doi.org/10.1016/j.sciaf.2024.e02281>
- Nguyen, P. T., Kend, M., & Le, D. Q. (2024). Digital transformation in Vietnam: The impacts on external auditors and their practices. *Pacific Accounting Review*, 36(1), 144–160. <https://doi.org/10.1108/par-04-2023-0051>
- Oldeweme, A., Märtins, J., Westmattellmann, D., & Schewe, G. (2021). The Role of Transparency, Trust, and Social Influence on Uncertainty Reduction in Times of Pandemics: Empirical Study on the Adoption of COVID-19 Tracing Apps. *Journal of Medical Internet Research*, 23(2), e25893. <https://doi.org/10.2196/25893>
- Pemer, F. (2021). Enacting Professional Service Work in Times of Digitalization and Potential Disruption. *Journal of Service Research*, 24(2), 249–268. <https://doi.org/10.1177/1094670520916801>
- Peters, C. P. H. (2022). Auditor Automation Usage and Professional Skepticism. *Auditor Automation Usage and Professional Skepticism*.
- Petty, R. E., & Cacioppo, J. T. (1986). *THE ELABORATION LIKELIHOOD MODEL OF PERSUASION*.
- Putri, J. W. (2024). Indonesia and ASEAN chairmanship in 2023: Leading the region in strengthening relations with China. *International Journal of Law and Politics Studies*, 6(1), 96–106. <https://doi.org/10.32996/ijlps.2024.6.1.11>
- Rajak, M., & Shaw, K. (2021). An extension of technology acceptance model for mHealth user adoption. *Technology in Society*, 67, 101800. <https://doi.org/10.1016/j.techsoc.2021.101800>
- Rana, N. P., Pillai, R., Sivathanu, B., & Malik, N. (2024). Assessing the nexus of Generative AI adoption, ethical considerations and organizational performance. *Technovation*, 135, 103064. <https://doi.org/10.1016/j.technovation.2024.103064>

- Rathnayake, A. S., Nguyen, T. D. H. N., & Ahn, Y. (2025). Factors Influencing AI Chatbot Adoption in Government Administration: A Case Study of Sri Lanka's Digital Government. *Administrative Sciences*, 15(5), 157. <https://doi.org/10.3390/admsci15050157>
- Robinson, S. N., Curtis, M. B., & Robertson, J. C. (2018). Disentangling the Trait and State Components of Professional Skepticism: Specifying a Process for State Scale Development. *Auditing: A Journal of Practice & Theory*, 37(1), 215–235. <https://doi.org/10.2308/ajpt-51738>
- Schnackenberg, A. K., & Tomlinson, E. C. (2016). Organizational Transparency: A New Perspective on Managing Trust in Organization-Stakeholder Relationships. *Journal of Management*, 42(7), 1784–1810. <https://doi.org/10.1177/0149206314525202>
- Seethamraju, R., & Hecimovic, A. (2023). Adoption of artificial intelligence in auditing: An exploratory study. *Australian Journal of Management*, 48(4), 780–800. <https://doi.org/10.1177/03128962221108440>
- Shin, D. (2020). User perceptions of algorithmic decisions in the personalized AI system: perceptual Evaluation of fairness, accountability, transparency, and explainability. *Journal of Broadcasting & Electronic Media*, 64(4), 541–565. <https://doi.org/10.1080/08838151.2020.1843357>
- Shin, D. (2021). The effects of explainability and causability on perception, trust, and acceptance: Implications for explainable AI. *International Journal of Human-Computer Studies*, 146, 102551. <https://doi.org/10.1016/j.ijhcs.2020.102551>
- Shin, D., & Park, Y. J. (2019). Role of fairness, accountability, and transparency in algorithmic affordance. *Computers in Human Behavior*, 98, 277–284. <https://doi.org/10.1016/j.chb.2019.04.019>
- Statista. (2025a). *Artificial intelligence (AI) adoption, risks, and challenges*. <https://www.statista.com/study/132103/artificial-intelligence-ai-adoption-risks-and-challenges/>
- Statista. (2025b). *Indonesia: GenAI market size 2030*. Statista. <https://www-statista-com.unpad.idm.oclc.org/forecasts/1493407/indonesia-genai-market-size>
- Sullivan, V., & Weger, K. (2025). *Transparency and Explainability in AI-Assisted Decision Making: Effects on Trust, Perceived Reliability, Confidence, and Ease of Understanding*. 69(1).
- Surahman, R. H., & Legowo, N. (2024). Factors that affecting firm performance mediated by digital transformation in Indonesia's state-owned enterprises plantation and forestry industries. *Edelweiss Applied Science and Technology*, 8(6). <https://doi.org/10.55214/25768484.v8i6.2810>
- Tian, Y., Zhang, H., Jiang, Y., & Yang, Y. (2021). Understanding trust and perceived risk in sharing accommodation: An extended elaboration likelihood model and moderated by risk attitude. *Journal of Hospitality Marketing & Management*, 31(3), 348–368. <https://doi.org/10.1080/19368623.2022.1986190>
- Vitali, S., & Giuliani, M. (2024). Emerging digital technologies and auditing firms: Opportunities and challenges. *International Journal of Accounting Information Systems*, 53, 100676. <https://doi.org/10.1016/j.accinf.2024.100676>
- von Eschenbach, W. J. (2021). Transparency and the Black Box Problem: Why We Do Not Trust AI. *Philosophy & Technology*, 34(4), 1607–1622. <https://doi.org/10.1007/s13347-021-00477-0>
- Vössing, M., Köhl, N., Lind, M., & Satzger, G. (2022). Designing transparency for effective human-ai collaboration. *Information Systems Frontiers*, 24(3), 877–895. <https://doi.org/10.1007/s10796-022-10284-3>

- Wang, W., & Benbasat, I. (2016). Empirical Assessment of Alternative Designs for Enhancing Different Types of Trusting Beliefs in Online Recommendation Agents. *Journal of Management Information Systems*, 33(3), 744–775. <https://doi.org/10.1080/07421222.2016.1243949>
- Yu, L., & Li, Y. (2022). Artificial Intelligence Decision-Making Transparency and Employees' Trust: The Parallel Multiple Mediating Effect of Effectiveness and Discomfort. *Behavioral Sciences*, 12(5), 127. <https://doi.org/10.3390/bs12050127>
- Zhang, T., Tao, D., Qu, X., Zhang, X., Zeng, J., Zhu, H., & Zhu, H. (2020). Automated vehicle acceptance in China: Social influence and initial trust are key determinants. *Transportation Research Part C: Emerging Technologies*, 112, 220–233. <https://doi.org/10.1016/j.trc.2020.01.027>
- Zhao, W. X., Zhou, K., Li, J., Tang, T., Wang, X., Hou, Y., Min, Y., Zhang, B., Zhang, J., Dong, Z., Du, Y., Yang, C., Chen, Y., Chen, Z., Jiang, J., Ren, R., Li, Y., Tang, X., Liu, Z., ... Wen, J.-R. (2023). A survey of large language models. <https://arxiv.org/abs/2303.18223>
- Zhou, T. (2012). Understanding users' initial trust in mobile banking: An elaboration likelihood perspective. *Computers in Human Behavior*, 28(4), 1518–1525. <https://doi.org/10.1016/j.chb.2012.03.021>
- Zhou, T., & Wu, X. (2025). Examining generative AI user disclosure intention: An ELM perspective. *Universal Access in the Information Society*, 24(2), 1209–1220. <https://doi.org/10.1007/s10209-024-01130-1>