

AI Content Provenance Governance and Managerial Decision Quality: Evidence from Knowledge-Intensive Service Firms

Weiwei Liu

Jiujiang Polytechnic University of Science and Technology
No.1 Gongqing Avenue, Gongqingcheng City, Jiujiang City, Jiangxi Province, 332020, P.R.China
3636576772@qq.com

Received date: July 6, 2026, revision date: July 28, 2026, Accepted: August 30, 2026

ABSTRACT

This study examines whether AI content provenance governance is associated with managerial decision quality in knowledge-intensive service firms and whether evidence verification accounts for this association. Drawing on organizational information processing theory, a survey of 412 managers across management consulting, software and IT services, and engineering consulting firms was conducted. Managers assessed AI content governance arrangements in their most recent project decision, reported verification behaviors, and completed a scenario-based AI literacy test. Decision quality was independently rated by two trained evaluators on four dimensions. The results show that AI content provenance governance has a significant positive association with managerial decision quality (coefficient = 3.71, $p < .001$). This association operates partly through evidence verification, with a positive indirect effect at the sample mean level of AI literacy (estimate = 1.81, 95% CI [0.94, 2.72]). Evidence verification is positively associated with decision quality after accounting for governance and controls (coefficient = 0.38, $p < .001$). The hypothesized positive moderation of AI literacy on the governance-verification link does not reach conventional significance ($p = .051$), although conditional slope analysis indicates a stronger governance-verification association at higher literacy levels. These findings contribute to digital management research by distinguishing provenance records, verification behaviors, and decision outcomes as separate constructs, and suggest that governance arrangements may improve decisions when they support rather than replace evidence checking.

Keywords: AI Content Provenance Governance, Evidence Verification, Managerial Decision Quality, AI Literacy, Knowledge-Intensive Services, Digital Management

1. Introduction

Managers in knowledge-intensive service firms increasingly rely on AI-generated analyses, recommendations, and project materials when making decisions for clients (Brynjolfsson et al., 2025; Noy and Zhang, 2023). This reliance creates a practical challenge: AI-generated content can include fabricated references, unsupported claims, and recommendations that appear coherent yet contradict relevant evidence (Ji et al., 2023; Walters and Wilder, 2023). A systematic review and meta-analysis of human-AI combinations found that combined performance does not consistently exceed the better-performing party, and that outcomes depend substantially on task characteristics and context (Vaccaro et al., 2024). These findings suggest that organizational arrangements governing how AI content enters managerial decisions deserve closer attention, particularly regarding whether such arrangements help managers evaluate the evidence underlying AI-generated materials.

Within the field of digital management and organizational information processing, this study investigates whether AI content provenance governance, defined by the availability of source records, version tracking,

applicability documentation, and review responsibilities for AI-generated project materials, is associated with higher managerial decision quality, and through what behavioral mechanism this association may operate. Organizational information processing theory posits that when task uncertainty exceeds existing routines, organizations benefit from arrangements that enhance their capacity to acquire and process decision-relevant information (Galbraith, 1974). Applied to AI-assisted decisions, provenance governance represents a set of organizational practices that may supply managers with the traceable information needed to assess what AI content is based on, how it was generated, and who is responsible for reviewing it (Mäntymäki et al., 2022; Birkstedt et al., 2023; Papagiannidis et al., 2025).

Yet the presence of governance records does not by itself ensure that managers will consult, interpret, and act on the evidence these records make accessible. Raisch and Krakowski (2021) observe that AI adoption in management involves a persistent tension between automating tasks for efficiency and augmenting human judgment. The availability of source documentation may lower the cost of evidence checking, but whether managers actually perform such checks, and whether checking leads to better decisions, remains an open question. Current research on AI governance has largely focused on principles, frameworks, and institutional arrangements at the organizational level (Birkstedt et al., 2023; Papagiannidis et al., 2025), without tracing how specific governance practices connect through managerial behavior to decision outcomes.

This study addresses three research questions in the context of knowledge-intensive service firms:

RQ1: Is AI content provenance governance positively associated with managerial decision quality?

RQ2: Does evidence verification account for a statistical indirect association between governance and decision quality?

RQ3: Does managerial AI literacy moderate the association between governance and evidence verification?

To address these questions, this study surveys 412 managers across three knowledge-intensive service industries, namely management consulting, software and IT services, and engineering and technical consulting, who used AI-generated content in a recent project decision. A four-variable framework specifies AI content provenance governance as the independent variable, evidence verification behaviors as the mediator, managerial decision quality evaluated by independent raters as the dependent variable, and managerial AI literacy measured through scenario-based knowledge questions as the moderator. This separation responds to a recurring concern in AI governance research that records of provenance, actual checking of evidence, and decision outcomes are conceptually distinct yet frequently conflated in measurement (Enholm et al., 2022; Mikalef and Gupta, 2021).

The study makes three contributions. First, it operationalizes AI content provenance governance as a formative construct covering source linkage, AI contribution tracking, generation context, and review responsibility, thereby moving beyond general governance principles toward specific, assessable practices. Second, it models evidence verification as a distinct behavioral link between governance information and decision quality, without assuming that governance directly improves outcomes. Third, it examines AI literacy as a potential boundary condition that may shape how managers use governance information, while accounting for two competing possibilities: literacy may complement governance by enabling more effective use of provenance records, or it may substitute for governance if knowledgeable managers verify evidence regardless of available documentation.

2. Literature Review and Hypothesis Development

2.1 Organizational Information Processing and Construct Boundaries

Galbraith (1974) proposed that organizations face uncertainty when the information required for a decision exceeds what is available through existing structures. Organizations can respond by reducing

information processing demands or by increasing their capacity to acquire, transmit, and interpret relevant information. This framework provides a theoretical basis for understanding how AI content provenance governance may function as an organizational arrangement that increases the information available to managers who use AI-generated materials for project decisions.

In the context of AI-assisted managerial work, information processing demands arise because AI-generated content, including analyses, summaries, and recommendations, may contain errors, fabrications, or context-dependent claims that are not immediately apparent (Ji et al., 2023; Walters and Wilder, 2023). Organizational AI governance encompasses the rules, practices, processes, and technical tools through which organizations direct their use of AI in alignment with strategic and regulatory requirements (Mäntymäki et al., 2022). Within this broader domain, this study concentrates on content provenance governance: the organizational arrangements that make source origins, processing history, version records, and review responsibilities traceable and accessible for specific AI-generated materials. The provenance concept itself has roots in data management research, where Herschel et al. (2017) describe provenance as the record of where data originated, how it was processed, and what transformations it underwent. Recent governance research extends these ideas to organizational settings, identifying coordination mechanisms, responsibility structures, and process-level accountability as key elements (Papagiannidis et al., 2025).

The distinction between provenance information and verification behavior is important. Provenance records make evidence checking possible, but they do not guarantee that checking occurs. The AI governance literature has documented a gap between governance principles and actual implementation in organizations (Birkstedt et al., 2023). The present study addresses this gap by separately measuring: (1) the availability of governance arrangements, (2) the extent of evidence verification behavior, (3) the quality of the resulting decision, and (4) the manager's AI-related knowledge.

2.2 AI Content Provenance Governance and Managerial Decision Quality

When managers make project decisions using AI-generated materials, the quality of those decisions depends in part on the quality of information available to them (Galbraith, 1974). AI content provenance governance may improve this information quality by enabling managers to identify where content originated, whether it has been modified, what its known limitations are, and who is accountable for reviewing it.

Several related lines of evidence support a positive association. Research on AI and business value indicates that the organizational conditions surrounding AI use, instead of AI capabilities alone, shape whether AI investments translate into performance improvements (Enholm et al., 2022). Case-based evidence from AI governance implementation identifies coordination, responsibility assignment, and information documentation as practices that may support more effective AI use in organizations (Papagiannidis et al., 2023). Structured documentation of data origins, intended use conditions, and known limitations provides users with contextual information that may sharpen their ability to assess whether AI outputs are appropriate for specific tasks (Geburu et al., 2021). Design principles for AI transparency further argue that traceability should be embedded in system development rather than appended afterward, so that provenance information is available at the point of decision rather than buried in technical logs (Felzmann et al., 2020).

At the same time, governance records do not automatically improve decisions. Documentation that is voluminous but poorly structured may increase cognitive load without aiding judgment (Mikalef and Gupta, 2021). Records may also be accepted at face value without genuine scrutiny, and documentation itself may contain errors. These considerations suggest that the governance-quality association, if it exists, likely depends on whether managers engage with the available information. As a starting point, this study first examines the overall statistical association:

H1: AI content provenance governance is positively associated with managerial decision quality.

2.3 Governance and Evidence Verification

A direct governance-quality association, even if observed, leaves the behavioral mechanism unspecified. This study proposes that evidence verification, defined as the actual checking of source existence, claim support, contextual applicability, and independent corroboration, serves as an intermediate behavioral link. These verification activities require both effort and access to traceable records.

Governance arrangements may facilitate verification by reducing the costs of search and comparison. When source links, version histories, and applicability notes are readily available, the cost of checking whether a claim is supported by its cited source decreases. Research on AI-assisted decision-making suggests that when the cost of independent evaluation falls, people are more likely to engage in effortful checking rather than accepting AI outputs passively (Vasconcelos et al., 2023). In professional medical settings, practitioners who actively explored the relationship between AI outputs and clinical evidence were better positioned to integrate AI into their critical judgments (Lebovitz et al., 2022). An evidence base for AI verification can also be understood in terms of citation and reference integrity: AI systems can generate seemingly complete bibliographic entries that turn out to be fabricated or that misattribute claims to genuine sources (Ji et al., 2023; Walters and Wilder, 2023), making source-level checking a concrete verification task.

Accessibility of records, however, does not guarantee their use. Governance records may exist but remain unconsulted if managers perceive checking as unnecessary or lack time. The distinction between record availability and verification behavior parallels findings that interpretability tools for machine learning models may go unused or be misinterpreted even by technically skilled users (Fügner et al., 2022). Design guidelines for trustworthy AI systems emphasize that verifiability, defined as the capacity for users to confirm system outputs against independent evidence, should be treated as a distinct design objective alongside transparency (Shneiderman, 2020). These possibilities motivate empirical examination.

H2: AI content provenance governance is positively associated with evidence verification.

2.4 Evidence Verification and the Indirect Association

If governance supports verification (H2), and if verification improves decision quality, then evidence verification may account for part of the association between governance and decision quality. This indirect association represents a statistical decomposition rather than a causal claim.

Evidence verification, as operationalized in this study, encompasses four activities: checking whether cited sources actually exist, whether sources support the specific claims attributed to them, whether evidence remains applicable to the current decision context, and whether independent sources corroborate key judgments. Each activity involves comparing AI-generated claims against external reference points. In professional settings, practitioners who invest more effort in verifying AI suggestions against their own knowledge and available evidence show different decision patterns than those who accept or reject suggestions without independent checking (Jussupow et al., 2021; Lebovitz et al., 2022).

Experimental evidence further supports the role of effortful engagement. Cognitive forcing functions that require users to form independent judgments before seeing AI recommendations reduce overreliance on incorrect suggestions, though at the cost of additional effort and lower subjective satisfaction (Buçinca et al., 2021). This finding implies that the benefits of verification come with costs, and that managers' willingness to bear these costs may depend on how governance arrangements structure the checking process.

The trust literature provides additional context for understanding why verification matters separately from trust. Lee and See (2004) distinguish between trust, reliance, and appropriate use of automated systems, noting that trust calibration is associated with better outcomes only when it aligns with actual system capability. Glikson and Woolley (2020) review how transparency, reliability, and interaction context shape human trust in AI, but note that the relationship between trust and decision outcomes is not straightforward. Higher trust may lead to less checking, or it may accompany more informed engagement depending on whether trust is based on understanding or on deference. For this reason, the present study measures verification behavior directly rather than treating trust as a proxy for decision quality.

The indirect association is tested using the product of coefficients approach with bootstrap confidence intervals, following current recommendations for mediation analysis (Hayes, 2009). This approach does not establish causation but estimates the magnitude and uncertainty of the statistical indirect pathway. The cross-sectional design means that temporal ordering is assumed rather than observed, and unmeasured confounders may influence both verification and decision quality.

H3: Evidence verification is positively associated with managerial decision quality, conditional on AI content provenance governance and the prespecified controls.

H4: Evidence verification accounts for a positive statistical indirect association between AI content provenance governance and managerial decision quality.

2.5 Managerial AI Literacy as a Boundary Condition

Not all managers may benefit equally from governance arrangements. The extent to which a manager can interpret provenance records, recognize when claims lack adequate support, and evaluate the applicability of AI-generated content likely depends on their understanding of how AI systems work, what kinds of errors they produce, and how evidence quality should be assessed. This pre-existing knowledge, termed managerial AI literacy, may moderate the governance-verification association.

AI literacy has been conceptualized as a multi-dimensional construct encompassing understanding of AI capabilities and limitations, the ability to evaluate AI outputs, and the capacity to interact with AI systems in context (Ng et al., 2021). Measurement approaches range from self-report scales assessing perceived competence across multiple dimensions (Carolus et al., 2023; Laupichler et al., 2023; Wang et al., 2023) to objective knowledge tests measuring factual understanding (Hornberger et al., 2023). A systematic review of AI literacy scales finds that while internal consistency has been established for several instruments, content validity and cross-population applicability remain areas requiring further evidence (Lintner, 2024). Variation in how the construct is operationalized and what outcomes it predicts is also documented in a comprehensive review of AI literacy research (Pinski and Benlian, 2024).

The present study adopts an objective, scenario-based measurement approach. Managers respond to ten questions that assess their knowledge about the distinction between language fluency and factual accuracy, the limits of traceability as a correctness indicator, the need for independent verification of AI-generated references, the difference between source existence and claim support, the task-specificity of AI performance, and related judgment scenarios. This approach directly measures knowledge relevant to evidence evaluation rather than self-assessed confidence or general AI awareness.

Two competing mechanisms suggest different moderating directions. A complementarity argument holds that managers with greater AI literacy can extract more value from governance records because they better understand what to look for and how to interpret provenance information, thereby strengthening the governance-verification link. A substitution argument suggests that highly literate managers may verify

AI content regardless of governance arrangements, relying on their own knowledge to identify potential problems; under this pattern, the difference in verification between high-governance and low-governance conditions would be smaller for literate managers. Evidence from related contexts supports both possibilities. Brynjolfsson et al. (2025) find that AI productivity effects are more pronounced for less experienced workers, consistent with substitution, while Jia et al. (2024) find that AI augments creativity most for employees with relevant skills, consistent with complementarity. Krakowski et al. (2023) note that the relationship between traditional expertise and AI-augmented performance is not fixed but depends on the competitive and technological context.

Given these competing expectations, this study retains the positive moderation hypothesis as the primary prediction while acknowledging the substitution pattern as a plausible alternative that the data may support instead:

H5: Managerial AI literacy positively moderates the association between AI content provenance governance and evidence verification.

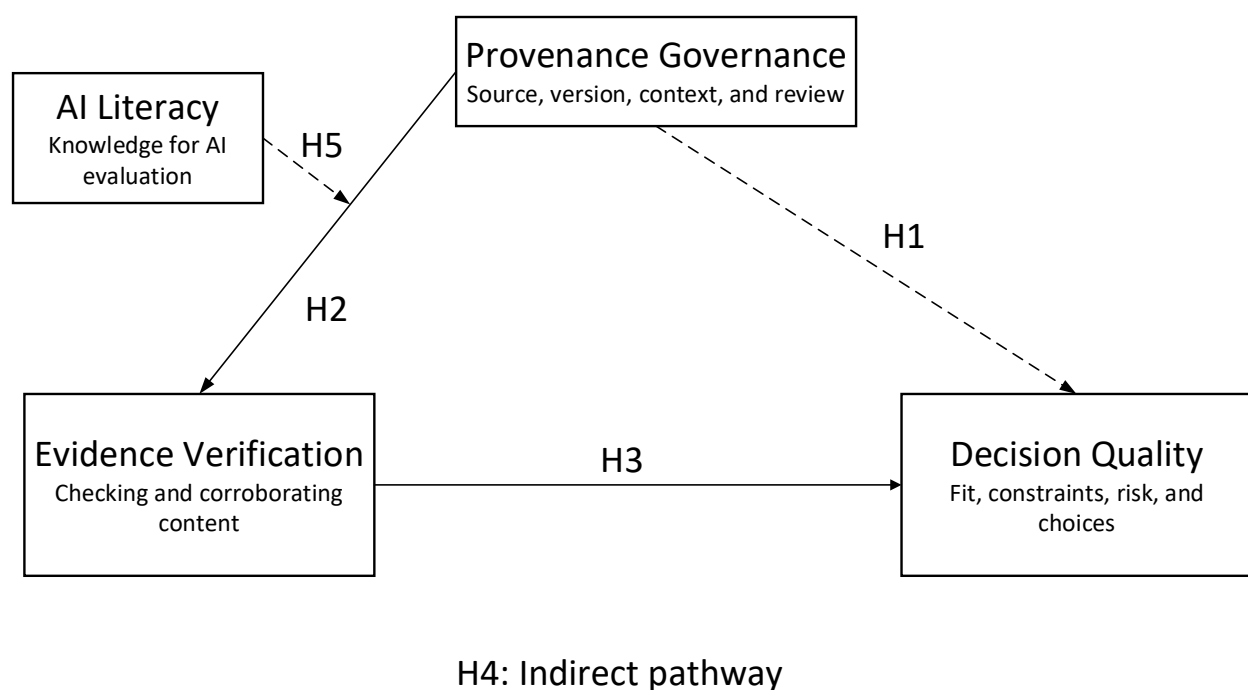


Figure 1: Conceptual Framework

3. Research Methodology

3.1 Research Design and Decision Context

This study adopts a cross-sectional survey design targeting managers in knowledge-intensive service firms who had used AI-generated content in a project decision within the preceding six months. Three industries were selected to represent distinct knowledge-intensive service contexts: management consulting, software and IT services, and engineering and technical consulting. Each industry involves different combinations of analytical complexity, client interaction, and technical judgment, providing variation in the conditions under which AI content governance operates.

The unit of analysis is a single decision episode in which the manager used AI-generated materials, such as analyses, summaries, or recommendations, as input to a project-level decision. Respondents assessed the governance arrangements available during that episode, reported the extent to which they verified the evidence underlying the AI content, and completed a scenario-based AI literacy test. Decision quality was evaluated separately by two trained raters who scored the decision outcomes on standardized rubric dimensions. Seven control variables were measured before the focal constructs: managerial experience in years, prior AI use experience in months, firm size (natural log of employees), pre-existing digital maturity, task complexity, and two industry indicator variables with management consulting as the reference category.

The study does not manipulate governance conditions or randomly assign managers to treatments. The statistical associations reported reflect observed patterns in a specific sample; causal claims require designs that address selection into governance conditions and temporal ordering.

3.2 Sample Construction and Descriptive Profile

Online questionnaires were distributed through professional networks and industry associations to managers in the three target industries, with a target of 150 responses per industry. Data collection took place between October and December 2025. A total of 450 completed responses were obtained. Each response was screened against prespecified quality rules: respondents who failed both embedded attention-check items and completed the questionnaire in fewer than 150 seconds were excluded, as were respondents for whom at least one of the four construct scores could not be computed due to missing item responses. These criteria led to the exclusion of 38 responses, yielding an analysis sample of 412 valid decision episodes. The exclusions included 7 cases meeting both the attention and time thresholds and 32 cases with insufficient item coverage for at least one construct; some cases met both criteria. The screening rules were defined before examining any focal associations. Across the 412 retained episodes, 1.81% of individual measurement cells were originally blank and were retained as missing, with construct scores computed from available items following construct-specific coverage rules described below.

Table 1 reports the descriptive statistics and bivariate correlations for all study variables.

Table 1: Descriptive Statistics and Correlations														
No.	Variable	Mean	SD	1	2	3	4	5	6	7	8	9	10	11
1	AI Content Provision	4.09	1.13	1.00										
2	Governance Evidence Verification	52.76	24.11	0.27	1.00									
3	Managerial Decision	51.68	22.04	0.22	0.49	1.00								

No.	Variable	Mean	SD	1	2	3	4	5	6	7	8	9	10	11
4	sion Quality Managerial AI Lite racy Managerial Exp erie nce Prio r AI Use Exp erie nce Fir m Size (Nat ural Log of Emp loye es)	52.7 2	21.6 4	0.03	0.11	0.14	1.00							
5	Pre- Exis ting Digi tal Mat urity Tas k Co mpl exit y Soft war e and IT Serv ices Engi neer ing and	7.39	4.84	- 0.03	- 0.06	0.05	0.10	1.00						
6		24.4 0	10.4 1	- 0.03	0.00	0.02	0.09	- 0.01	1.00					
7		4.52	0.97	0.18	0.05	0.03	0.07	- 0.02	- 0.01	1.00				
8		4.07	1.80	0.30	0.22	0.15	0.14	- 0.11	0.10	0.20	1.00			
9		3.71	1.88	- 0.01	- 0.20	- 0.27	- 0.06	0.08	- 0.01	0.20	- 0.06	1.00		
10		0.33	0.47	0.07	0.09	0.08	0.04	0.01	- 0.02	0.08	0.24	0.00	1.00	
11		0.33	0.47	- 0.02	0.01	- 0.03	- 0.05	0.00	- 0.04	- 0.11	- 0.15	- 0.02	- 0.50	1.00

No.	Variable	Mean	SD	1	2	3	4	5	6	7	8	9	10	11
	Tech nical Con sulti ng													

Note: N = 412. AI Content Provenance Governance uses a 1 to 7 scale; Evidence Verification, Managerial Decision Quality, and Managerial AI Literacy use 0 to 100 scales. Managerial Experience is measured in years; Prior AI Use Experience in months. Industry indicators are coded as 0 or 1 with management consulting as the reference category.

The sample is evenly distributed across the three industries. Governance scores span the full scale range with a mean near the midpoint ($M = 4.09$, $SD = 1.13$). Evidence verification ($M = 52.76$, $SD = 24.11$) and decision quality ($M = 51.68$, $SD = 22.04$) both show substantial variation. AI literacy has a comparable distribution ($M = 52.72$, $SD = 21.64$). Among the correlations, verification and decision quality show the strongest bivariate association ($r = 0.49$), followed by governance-verification ($r = 0.27$) and governance-quality ($r = 0.22$). AI literacy shows modest correlations with all other focal variables ($r = 0.03$ to 0.14). Task complexity is negatively associated with both verification ($r = -0.20$) and decision quality ($r = -0.27$), consistent with the expectation that more complex tasks are associated with lower performance indicators.

3.3 Measures and Scoring

Table 2 summarizes the construct operationalization, measurement structure, and scoring approach for each variable. Twenty-six items yield 30 raw measurement fields because the decision quality construct retains two independent ratings per dimension. All items were developed for this study based on conceptual sources identified in the literature review; no items were taken verbatim from previously validated scales, and the measurement properties reported here apply to this sample only.

Table 2: Construct Operationalization and Scoring

Construct	Measurement Structure	Items	Original Response Scale	Composite Scoring	Conceptual Sources
AI Content Provenance Governance	Four formative components	8	1 to 7	Mean within each two-item component; equal mean of four components; at least 7 items observed and every component represented	Mäntymäki et al. (2022); Herschel et al. (2017); Felzmann et al. (2020); Mökander et al. (2021)
Evidence Verification	Four behavioral indicators	4	0 to 4	Mean of all four observed indicators multiplied by 25	Walters and Wilder (2023); Buçinca et al. (2021); Lebovitz et al. (2022)
Managerial Decision Quality	Four rubric dimensions with two independent ratings each	4 dimensions; 8 rating fields	0 to 4	Mean available ratings within each dimension; mean of all	Jussupow et al. (2021); Vaccaro et al. (2024), conceptual

Construct	Measurement Structure	Items	Original Response Scale	Composite Scoring	Conceptual Sources
				four dimensions multiplied by 25	support only
Managerial AI Literacy	Scenario-based knowledge test	10	A to D; E indicates do not know	Number correct multiplied by 10; wrong, E, and blank score zero; at least 8 attempted items required	Ng et al. (2021); Hornberger et al. (2023); Carolus et al. (2023)

AI Content Provenance Governance comprises four formative components, each measured by two items on a 1 to 7 scale: source linkage (whether AI claims have identifiable original sources and version-stamped records), AI contribution and versioning (whether AI-generated and human-revised portions are distinguishable with tracked revisions), generation context and applicability boundaries (whether model versions, input conditions, and known limitations are documented), and review responsibility and error correction (whether review assignments and correction procedures are specified). Items evaluate the availability of these governance arrangements for the specific decision episode, not the respondent's general perception of organizational policy. The composite score is the equal-weighted mean of the four component means.

Evidence Verification consists of four behavioral indicators scored from 0 (not performed) to 4 (performed on more than 75% of key materials): source existence checking, claim-source support checking, temporal and contextual applicability checking, and independent corroboration. The composite is the mean of available indicators multiplied by 25, yielding a 0 to 100 scale.

Managerial Decision Quality is assessed by two independent raters on four rubric dimensions: objective fit, constraint compliance, risk treatment, and choice consistency. Each dimension uses a 0 to 4 anchor scale. The composite averages available ratings within each dimension and then averages across dimensions, multiplied by 25. This independent rating approach separates decision outcome measurement from the respondent's self-reported governance and verification behaviors.

Managerial AI Literacy uses ten scenario-based knowledge questions with four substantive options and a "do not know" option. Each correct answer scores 1; wrong answers, "do not know" responses, and blanks score 0. The total is computed over all ten items and multiplied by 10, yielding a 0 to 100 scale. Respondents who attempted fewer than eight items were excluded during quality screening.

3.4 Analytical Strategy and Uncertainty

All models are estimated using ordinary least squares (OLS) with heteroskedasticity-consistent standard errors (HC3) and two-sided t-tests. Governance and AI literacy scores are mean-centered within the analysis sample before entering the models. Two industry indicator variables are included with management consulting as the reference. Three equations are estimated:

$$\text{Overall association: } Y = \beta_0 + cX_c + dW_c + \gamma C + \varepsilon$$

$$\text{Verification model: } M = \alpha_0 + a_1X_c + a_2W_c + a_3(X_c \times W_c) + \gamma C + \varepsilon_M$$

$$\text{Decision quality model: } Y = \theta_0 + c'X_c + bM + dW_c + \gamma C + \varepsilon_Y$$

The indirect association at a given centered AI literacy level w is computed as $(a_1 + a_3w)b$. The index of moderated mediation equals a_3b . Bootstrap confidence intervals for the indirect association are based on 5,000 paired-equation resamples drawn with replacement from the same 412 records, reestimating both

the verification and decision quality models in each resample. Percentile intervals are reported. The analytical strategy follows established recommendations for rigorous method evaluation and adequate sample sizing in applied research (Lakens, 2022; Morris et al., 2019).

Robustness checks compare the primary specification against four alternatives: a strict complete-case analysis limited to respondents with no missing measurement fields, a specification that relaxes the speed and attention exclusion criteria, a specification that adds a direct governance-by-literacy interaction in the decision quality model, and an alternative AI literacy scoring rule that divides correct answers by attempted items rather than the fixed ten-item base. All specifications use the same model structure and HC3 inference.

4. Results

4.1 Measurement Quality

Table 3 reports measurement quality diagnostics for the analysis sample.

Table 3: Measurement Quality Diagnostics				
Measure	Diagnostic	Evaluation Sample	Value	Interpretation
AI Content Provenance Governance	Maximum component variance inflation factor	412 eligible episodes	1.51	Formative component collinearity diagnostic Limited internal consistency;
Managerial AI Literacy	KR-20	357 fully answered tests	0.55	instrument requires further development
All measurement fields	Missing item cells (%)	13,500 raw item cells	1.81	Original missing entries retained
Objective Fit	Exact agreement (%)	444 paired ratings	51.35	Inter-rater agreement
Objective Fit	Agreement within one point (%)	444 paired ratings	97.52	Inter-rater agreement
Objective Fit	Mean absolute rating difference	444 paired ratings	0.51	Inter-rater agreement
Constraint Compliance	Exact agreement (%)	439 paired ratings	53.76	Inter-rater agreement
Constraint Compliance	Agreement within one point (%)	439 paired ratings	96.36	Inter-rater agreement
Constraint Compliance	Mean absolute rating difference	439 paired ratings	0.50	Inter-rater agreement
Risk Treatment	Exact agreement (%)	443 paired ratings	56.88	Inter-rater agreement
Risk Treatment	Agreement within one point (%)	443 paired ratings	97.29	Inter-rater agreement
Risk Treatment	Mean absolute rating difference	443 paired ratings	0.46	Inter-rater agreement
Choice Consistency	Exact agreement (%)	443 paired ratings	53.50	Inter-rater agreement
Choice Consistency	Agreement within one point (%)	443 paired ratings	97.07	Inter-rater agreement
Choice Consistency	Mean absolute rating difference	443 paired ratings	0.49	Inter-rater agreement

The formative governance composite shows acceptable collinearity among its four components (maximum VIF = 1.51), indicating that the components contribute distinct information to the composite.

The AI literacy knowledge test has limited internal consistency ($KR-20 = 0.55$, based on 357 respondents who answered all ten items). This value is below conventional thresholds for well-established instruments and reflects the early stage of instrument development; the ten items span multiple knowledge domains, and the relatively modest consistency should be considered when interpreting results involving this variable. Inter-rater agreement on the decision quality dimensions is acceptable: exact agreement ranges from 51.35% to 56.88% across the four dimensions, and agreement within one scale point exceeds 96% for all dimensions. Mean absolute rating differences are approximately 0.50 scale points on the 0 to 4 scale, suggesting that the two raters produced broadly consistent evaluations.

4.2 Regression Estimates

Table 4 presents the results of the three regression models. All models use the same 412-case analysis sample, HC3 standard errors, and the same set of control variables.

Overall Association: $N = 412$, residual $df = 402$, $R^2 = 0.15$, adjusted $R^2 = 0.13$, HC3 Wald $F(9, 402) = 8.39$, $p < .001$.

Verification Model: $N = 412$, residual $df = 401$, $R^2 = 0.15$, adjusted $R^2 = 0.13$, HC3 Wald $F(10, 401) = 7.13$, $p < .001$.

Decision Quality Model: $N = 412$, residual $df = 401$, $R^2 = 0.30$, adjusted $R^2 = 0.28$, HC3 Wald $F(10, 401) = 20.11$, $p < .001$.

Table 4: Regression Estimates for Evidence Verification and Managerial Decision Quality

Model	Predictor	Coefficient	HC3 SE	t	p
Overall Association	Intercept	53.76	6.81	7.89	$p < .001$
Overall Association	AI Content Provenance Governance	3.71	0.96	3.85	$p < .001$
Overall Association	Managerial AI Literacy	0.10	0.05	2.00	$p = .047$
Overall Association	Managerial Experience	0.35	0.21	1.64	$p = .102$
Overall Association	Prior AI Use Experience	0.02	0.10	0.20	$p = .844$
Overall Association	Firm Size (Natural Log of Employees)	0.71	1.18	0.60	$p = .547$
Overall Association	Pre-Existing Digital Maturity	0.68	0.65	1.04	$p = .297$
Overall Association	Task Complexity	-3.21	0.56	-5.73	$p < .001$
Overall Association	Software and IT Services	2.22	2.49	0.89	$p = .374$
Overall Association	Engineering and Technical Consulting	0.14	2.54	0.05	$p = .957$
Verification Model	Intercept	52.43	6.90	7.60	$p < .001$
Verification Model	AI Content Provenance Governance	4.73	1.09	4.36	$p < .001$
Verification Model	Managerial AI Literacy	0.09	0.06	1.58	$p = .115$
Verification	AI Content	0.10	0.05	1.96	$p = .051$

Model	Predictor	Coefficient	HC3 SE	t	p
Model	Provenance Governance × Managerial AI Literacy				
Verification Model	Managerial Experience	-0.14	0.24	-0.60	p = .551
Verification Model	Prior AI Use Experience	0.00	0.10	-0.04	p = .970
Verification Model	Firm Size (Natural Log of Employees)	0.24	1.17	0.21	p = .836
Verification Model	Pre-Existing Digital Maturity	1.74	0.75	2.31	p = .021
Verification Model	Task Complexity	-2.44	0.62	-3.91	p < .001
Verification Model	Software and IT Services	3.36	2.82	1.19	p = .233
Verification Model	Engineering and Technical Consulting	3.44	2.79	1.23	p = .218
Decision Quality Model	Intercept	33.67	6.35	5.30	p < .001
Decision Quality Model	AI Content Provenance Governance	1.93	0.88	2.20	p = .029
Decision Quality Model	Evidence Verification	0.38	0.04	9.35	p < .001
Decision Quality Model	Managerial AI Literacy	0.06	0.04	1.35	p = .177
Decision Quality Model	Managerial Experience	0.41	0.19	2.19	p = .029
Decision Quality Model	Prior AI Use Experience	0.03	0.10	0.30	p = .761
Decision Quality Model	Firm Size (Natural Log of Employees)	0.54	1.08	0.50	p = .616
Decision Quality Model	Pre-Existing Digital Maturity	0.09	0.59	0.16	p = .876
Decision Quality Model	Task Complexity	-2.30	0.52	-4.44	p < .001
Decision Quality Model	Software and IT Services	0.72	2.35	0.31	p = .760
Decision Quality Model	Engineering and Technical Consulting	-1.22	2.33	-0.52	p = .602

The overall association model confirms that governance is positively associated with decision quality (coefficient = 3.71, SE = 0.96, $t = 3.85$, $p < .001$), supporting H1. Each one-point increase in governance on the 1 to 7 scale is associated with a 3.71-point increase in decision quality on the 0 to 100 scale, holding other variables constant. Task complexity shows a negative association (coefficient = -3.21, $p < .001$), while AI literacy reaches marginal significance (coefficient = 0.10, $p = .047$).

In the verification model, governance is positively associated with evidence verification at the mean level of AI literacy (coefficient = 4.73, SE = 1.09, $t = 4.36$, $p < .001$), supporting H2. The governance-by-literacy interaction term is positive but does not reach the conventional .05 threshold (coefficient = 0.10,

SE = 0.05, $t = 1.96$, $p = .051$). H5 is therefore not supported at conventional significance levels, though the direction is consistent with the complementarity prediction. Pre-existing digital maturity (coefficient = 1.74, $p = .021$) and task complexity (coefficient = -2.44, $p < .001$) are the only control variables with significant associations.

In the decision quality model, evidence verification has a strong positive association with decision quality (coefficient = 0.38, SE = 0.04, $t = 9.35$, $p < .001$), supporting H3. After introducing verification, the governance coefficient decreases from 3.71 to 1.93 ($p = .029$) but remains significant, indicating that governance retains a direct statistical association with decision quality beyond the verification pathway. Managerial experience also reaches significance in this model (coefficient = 0.41, $p = .029$). The decision quality model explains 30% of the variance (adjusted $R^2 = 0.28$), compared to 15% for the overall association and verification models.

4.3 Indirect Associations

Table 5 reports the conditional indirect associations and the index of moderated mediation, estimated using 5,000 paired-equation bootstrap resamples.

Table 5: Conditional Indirect Associations and the Index of Moderated Mediation

Association	AI Literacy Reference	Estimate	Bootstrap SE	95% Percentile CI	Resamples	N
Conditional indirect association	Mean (52.72)	1.81	0.45	[0.94, 2.72]	5,000	412
Conditional indirect association	16th percentile (30.00)	0.90	0.65	[-0.41, 2.16]	5,000	412
Conditional indirect association	Median (50.00)	1.70	0.46	[0.82, 2.61]	5,000	412
Conditional indirect association	84th percentile (80.00)	2.89	0.72	[1.54, 4.41]	5,000	412
Index of moderated mediation	Per ten-point increase in AI literacy	0.40	0.21	[0.02, 0.82]	5,000	412

At the sample mean level of AI literacy, the indirect association through evidence verification is positive and the confidence interval excludes zero (estimate = 1.81, 95% CI [0.94, 2.72]), supporting H4. The indirect association varies across literacy levels: it is not significant at the 16th percentile (estimate = 0.90, 95% CI [-0.41, 2.16]) but is significant at the median (estimate = 1.70, 95% CI [0.82, 2.61]) and 84th percentile (estimate = 2.89, 95% CI [1.54, 4.41]).

The index of moderated mediation, scaled per ten-point increase in AI literacy, is 0.40 (95% CI [0.02, 0.82]). This bootstrap interval is narrow and its lower bound sits close to zero, indicating limited precision. The index suggests a positive trend in which the indirect association through verification increases with AI literacy, but the marginal evidence warrants cautious interpretation.

4.4 Conditional Governance Associations

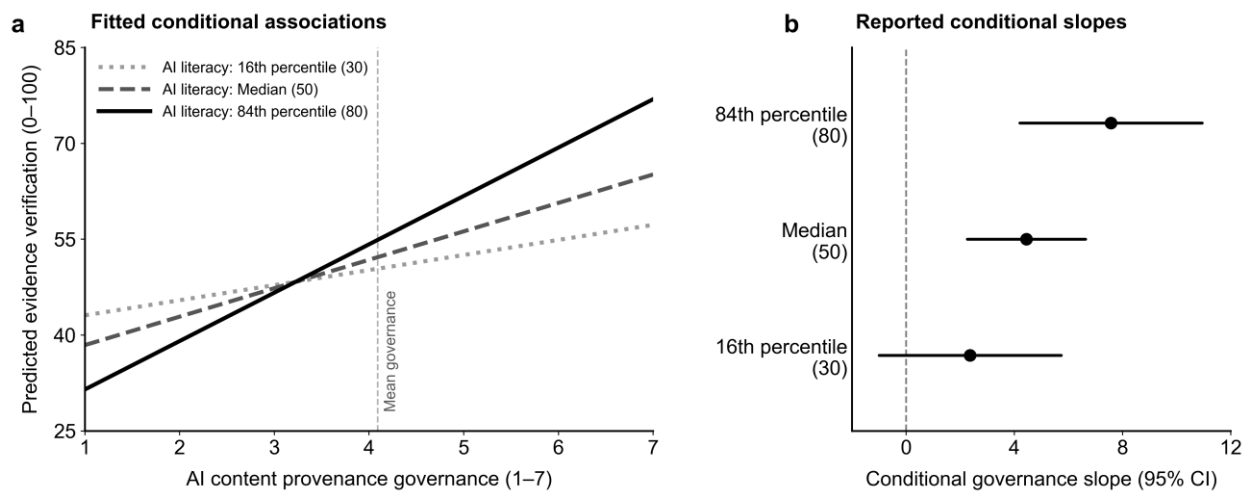
Table 6 and Figure 2 present the conditional associations between governance and evidence verification across the observed range of AI literacy.

Table 6: Conditional Associations Between AI Content Provenance Governance and Evidence Verification

AI Literacy Evaluation Point	Observed Score	Conditional Governance Slope	HC3 SE	t	p	95% CI
16th percentile	30.00	2.36	1.71	1.38	p = .168	[-1.00, 5.73]
Median	50.00	4.45	1.11	4.00	p < .001	[2.26, 6.63]
84th percentile	80.00	7.57	1.72	4.42	p < .001	[4.20, 10.94]

Note: Conditional slopes are derived from the verification model in Table 4. The slope at each evaluation point equals $a_1 + a_3w$, where w is the centered AI literacy score. Standard errors incorporate the HC3 covariance between a_1 and a_3 .

The governance-verification slope is not significant at the 16th percentile of AI literacy (slope = 2.36, $p = .168$) but is significant at the median (slope = 4.45, $p < .001$) and 84th percentile (slope = 7.57, $p < .001$). This pattern indicates that the governance-verification association strengthens with AI literacy, though the significance of individual conditional slopes does not by itself establish that moderation is present—the formal test of moderation is the interaction term reported in Table 4.

**Figure 2:** Conditional Governance Associations Across Managerial AI Literacy

4.5 Robustness Checks

Table 7 compares the focal coefficients across five analytical specifications.

Table 7: Robustness of the Regression Estimates

Specification	Focal Term and Model	N	Coefficient	HC3 SE	t	p
Primary scoring	Overall: AI Content Provenance Governance Verification:	412	3.71	0.96	3.85	p < .001
Primary scoring	AI Content Provenance Governance	412	4.73	1.09	4.36	p < .001

Specification	Focal Term and Model	N	Coefficient	HC3 SE	t	p
Primary scoring	Decision Quality: AI Content Provenance Governance	412	1.93	0.88	2.20	p = .029
Primary scoring	Decision Quality: Evidence Verification	412	0.38	0.04	9.35	p < .001
Primary scoring	Verification: Governance × AI Literacy	412	0.10	0.05	1.96	p = .051
Fully observed fields	Overall: AI Content Provenance Governance	268	4.27	1.22	3.50	p < .001
Fully observed fields	Verification: AI Content Provenance Governance	268	5.04	1.35	3.74	p < .001
Fully observed fields	Decision Quality: AI Content Provenance Governance	268	2.42	1.12	2.15	p = .032
Fully observed fields	Decision Quality: Evidence Verification	268	0.37	0.05	7.33	p < .001
Fully observed fields	Verification: Governance × AI Literacy	268	0.17	0.08	2.20	p = .029
No speed-attention exclusion	Overall: AI Content Provenance Governance	418	3.54	0.95	3.73	p < .001
No speed-attention exclusion	Verification: AI Content Provenance Governance	418	4.56	1.07	4.28	p < .001
No speed-attention exclusion	Decision Quality: AI Content Provenance Governance	418	1.82	0.86	2.12	p = .035
No speed-attention exclusion	Decision Quality: Evidence Verification	418	0.38	0.04	9.44	p < .001
No speed-attention exclusion	Verification: Governance × AI Literacy	418	0.11	0.05	2.11	p = .035

Specification	Focal Term and Model	N	Coefficient	HC3 SE	t	p
Direct moderation of quality	Overall: AI Content Provenance Governance	412	3.71	0.96	3.85	p < .001
Direct moderation of quality	Verification: AI Content Provenance Governance	412	4.73	1.09	4.36	p < .001
Direct moderation of quality	Decision Quality: AI Content Provenance Governance	412	1.88	0.89	2.12	p = .034
Direct moderation of quality	Decision Quality: Evidence Verification	412	0.39	0.04	9.33	p < .001
Direct moderation of quality	Verification: Governance × AI Literacy Decision Quality: Governance × AI Literacy	412	0.10	0.05	1.96	p = .051
Direct moderation of quality	Decision Quality: Governance × AI Literacy	412	-0.04	0.03	-1.17	p = .245
Knowledge per attempted item	Overall: AI Content Provenance Governance	412	3.68	0.96	3.82	p < .001
Knowledge per attempted item	Verification: AI Content Provenance Governance	412	4.72	1.09	4.35	p < .001
Knowledge per attempted item	Decision Quality: AI Content Provenance Governance	412	1.91	0.88	2.18	p = .030
Knowledge per attempted item	Decision Quality: Evidence Verification	412	0.38	0.04	9.32	p < .001
Knowledge per attempted item	Verification: Governance × AI Literacy	412	0.10	0.05	1.93	p = .055

The governance-quality, governance-verification, and verification-quality associations remain positive and significant across all five specifications. The main pattern, characterized by governance positively associated with verification, verification positively associated with decision quality, and a reduced but still significant direct governance effect, is stable.

The interaction term is more sensitive to analytical choices. In the primary analysis (N = 412), it falls just above the .05 threshold (p = .051). In the fully observed fields specification (N = 268), the interaction reaches significance (p = .029), as does the relaxed screening specification (N = 418, p = .035). The

alternative knowledge scoring rule produces $p = .055$. The direct moderation specification adds a governance-by-literacy interaction in the decision quality model, which is not significant (coefficient = -0.04 , $p = .245$), indicating no evidence of a direct moderating effect on decision quality. These results confirm that H1 through H4 are stable across reasonable analytical alternatives, while the evidence for H5 is sensitive to the sample definition and scoring rule, consistent with the marginal p value in the primary analysis.

5. Discussion

5.1 Interpretation and Comparison with Prior Research

This study set out to examine whether AI content provenance governance is associated with managerial decision quality in knowledge-intensive service firms, whether evidence verification accounts for this association, and whether managerial AI literacy moderates the governance-verification link. The results support four of the five hypotheses and provide a partially supported picture for the fifth.

The positive overall association between governance and decision quality (H1) is consistent with the broader organizational information processing perspective that decision quality improves when relevant information is available and accessible (Galbraith, 1974). The magnitude of the governance coefficient (3.71 points on the 0 to 100 decision quality scale per unit of governance on the 1 to 7 scale) represents a modest but meaningful association in practical terms, particularly given the range of organizational and individual variation captured by the control variables.

The verification pathway offers a more precise account of how governance relates to decisions. Evidence verification shows a strong association with decision quality (coefficient = 0.38 , $p < .001$), and the indirect association through verification is positive at the mean level of AI literacy (estimate = 1.81 , 95% CI [0.94 , 2.72]). This finding aligns with research showing that effortful engagement with AI outputs, instead of passive acceptance, leads to different decision patterns (Jussupow et al., 2021; Lebovitz et al., 2022). It is also consistent with experimental evidence that reducing the cost of independent checking increases the likelihood that users will identify errors in AI recommendations (Vasconcelos et al., 2023; Bućinca et al., 2021).

An important observation is that governance retains a significant direct association with decision quality (coefficient = 1.93 , $p = .029$) even after accounting for verification. This residual association may reflect governance functions not captured by the four verification indicators, for instance, the availability of review responsibility structures or error correction procedures may affect decision quality through channels other than the specific checking behaviors measured here. Alternatively, it may reflect unmeasured confounders associated with both governance and quality.

The marginal interaction result for H5 ($p = .051$) deserves careful interpretation. The direction is consistent with the complementarity prediction: higher AI literacy is associated with a stronger governance-verification link. Conditional slope analysis confirms that the governance-verification association is significant at the median and 84th percentile of literacy but not at the 16th percentile. This pattern suggests that managers with greater AI knowledge may be better equipped to translate governance records into effective verification behaviors, though the evidence is not strong enough to reject the null at conventional levels. The robustness analysis reveals that the interaction result varies with sample composition, reaching significance in the complete-case subsample ($p = .029$) but exceeding $.05$ in the primary and alternative scoring analyses. This sensitivity warrants caution against treating the moderation result as established.

The substitution pattern predicted by some prior work, in which AI benefits accrue primarily to less skilled or less experienced workers (Brynjolfsson et al., 2025), does not emerge clearly in these data, though it cannot be ruled out. The task contexts differ considerably: customer service interactions involve

structured, repetitive work, whereas the project decisions examined here involve more open-ended judgment under uncertainty. These differences may explain why the capability-AI interaction operates differently across settings, as others have also observed in studies of AI and creative work (Jia et al., 2024; Krakowski et al., 2023).

5.2 Theoretical Implications

This study contributes to research on AI governance and digital management in three ways.

First, it distinguishes governance information supply, verification behavior, and decision quality as separate constructs rather than treating governance as a direct determinant of performance. Much of the governance literature identifies principles, frameworks, and best practices (Birkstedt et al., 2023; Papagiannidis et al., 2025) without tracing behavioral pathways to decision outcomes. By positioning evidence verification as a mediating behavior, this study provides a process-level account: governance arrangements may improve decisions not directly, but by creating conditions under which managers are more likely to check the evidence behind AI-generated content. This process distinction matters because it implies that governance investments may be ineffective if they do not translate into actual verification activities.

Second, the operationalization of AI content provenance governance as a formative construct with four specific components, namely source linkage, AI contribution tracking, generation context, and review responsibility, moves beyond abstract governance principles toward assessable organizational practices. Each component corresponds to a specific type of information that a manager could use when evaluating AI content. This operationalization draws on data documentation practices (Gebru et al., 2021), provenance tracking research (Herschel et al., 2017), transparency-by-design principles (Felzmann et al., 2020), and ethics-based auditing frameworks (Mökander et al., 2021) to ground governance in concrete, observable arrangements.

Third, the competing moderation mechanisms of complementarity versus substitution provide a framework for understanding individual differences in AI governance effectiveness. Although the data do not conclusively distinguish these mechanisms, the pattern of results is more consistent with complementarity than substitution, suggesting that governance structures and individual knowledge may reinforce each other rather than serving as alternatives. This finding resonates with the automation-augmentation paradox discussed by Raisch and Krakowski (2021), where the relationship between human capability and AI support is neither purely additive nor purely substitutive but depends on the task and organizational context.

5.3 Practical Implications

The findings carry several implications for organizations that manage AI-assisted decision processes. These implications are framed as considerations informed by the observed associations rather than guaranteed prescriptions, given the cross-sectional design and the early stage of construct development.

First, organizations seeking to improve the quality of AI-assisted project decisions may benefit from investing in traceable provenance records rather than assuming that AI content quality alone determines decision outcomes. Governance practices that make source links, version histories, applicability boundaries, and review assignments accessible and actionable give managers the raw material for evidence checking. This recommendation aligns with research showing that AI governance effectiveness depends on operational implementation rather than stated principles alone (Papagiannidis et al., 2023). In service contexts where AI assists process-level decisions, establishing clear accountability structures for reviewing AI outputs has been identified as a practical priority (Kaškevičius, 2025).

Second, governance investments may be more effective when they are designed to support specific verification behaviors rather than to supply information in bulk. The four verification activities examined, namely checking source existence, claim support, contextual applicability, and independent corroboration, represent concrete steps that can be incorporated into project workflows. Organizations might consider building verification checklists into their AI content review procedures, particularly for decisions where the cost of error is high. The association between reduced verification costs and increased verification effort observed in experimental settings (Vasconcelos et al., 2023) suggests that making verification easier, for example through direct links to source materials rather than requiring independent searches, may be more productive than simply providing more documentation.

Third, the AI literacy results suggest that managers' ability to use governance information effectively may depend on their understanding of AI capabilities and limitations. While the moderation evidence is not definitive, the conditional slope pattern indicates that governance-verification associations are stronger for managers with higher AI literacy. This suggests potential value in training programs that develop managers' ability to evaluate AI-generated evidence, distinguish between source existence and source support, recognize the task-specificity of AI performance, and understand the difference between repeated AI outputs and independent verification. Such programs should focus on practical judgment scenarios relevant to the manager's decision context rather than on general AI awareness.

Fourth, the separation of governance, verification, and decision quality has measurement implications. Organizations that evaluate their AI governance effectiveness should consider measuring not just whether governance records exist, but whether managers actually use them and whether decision quality improves as a result. Relying solely on governance compliance measures, such as checking whether records are maintained, may overestimate the practical impact of governance if the behavioral link through verification is weak.

6. Conclusion

6.1 Main Conclusion

This study examined whether AI content provenance governance is associated with managerial decision quality in knowledge-intensive service firms, whether evidence verification mediates this association, and whether managerial AI literacy moderates the governance-verification relationship. Using survey data from 412 managers across management consulting, software and IT services, and engineering and technical consulting, the results show that governance is positively associated with both decision quality and evidence verification; that verification accounts for a positive indirect association between governance and decision quality; and that the hypothesized positive moderation by AI literacy falls just short of conventional significance. The key insight is that governance arrangements may improve decisions by facilitating evidence checking rather than by directly improving information quality, and that managers' AI-related knowledge may shape their ability to use governance information effectively.

6.2 Limitations and Future Research

Several limitations affect the interpretation of these findings. First, the cross-sectional design does not permit causal conclusions; the associations reported may reflect selection effects or unmeasured confounders rather than governance influencing verification and decision quality in the temporal order assumed by the model. Second, governance availability and verification behaviors are self-reported by the same respondent, raising the possibility of common method effects, though decision quality was independently rated and AI literacy was measured through objective knowledge questions. Third, the AI literacy knowledge test shows limited internal consistency ($KR-20 = 0.55$), reflecting the early stage of this instrument's development. Fourth, the interaction effect is sensitive to analytical specifications, indicating that the moderation result should be treated as preliminary. Fifth, some respondents had missing item

responses; the item-level missing data rate of 1.81% is low, but respondents with missing data may differ systematically from complete-case respondents.

Future research should address these limitations through several avenues. Longitudinal or experimental designs that manipulate governance conditions would provide stronger evidence for causal claims. Multi-source data collection, for example obtaining governance assessments from organizational records rather than respondent self-reports, would reduce common method concerns. The AI literacy instrument requires further development, including content validation with subject matter experts, test-retest reliability assessment, and evaluation across different managerial populations. The verification construct would benefit from behavioral observation or process tracking rather than retrospective self-report. Finally, extending this research to other service contexts and task types would test the generalizability of the governance-verification-quality pathway.

References

- Birkstedt, T., Minkkinen, M., Tandon, A. & Mäntymäki, M. (2023), AI Governance: Themes, Knowledge Gaps and Future Agendas, *Internet Research*, Vol. 33, No. 7, 133-167.
- Brynjolfsson, E., Li, D. & Raymond, L. (2025), Generative AI at Work, *The Quarterly Journal of Economics*, Vol. 140, No. 2, 889-942.
- Buçinca, Z., Malaya, M.B. & Gajos, K.Z. (2021), To Trust or to Think: Cognitive Forcing Functions Can Reduce Overreliance on AI in AI-Assisted Decision-Making, *Proceedings of the ACM on Human-Computer Interaction*, Vol. 5, No. CSCW1, Article 188, 1-21.
- Carolus, A., Koch, M.J., Straka, S., Latoschik, M.E. & Wienrich, C. (2023), MAILS - Meta AI Literacy Scale: Development and Testing of an AI Literacy Questionnaire Based on Well-Founded Competency Models and Psychological Change- and Meta-Competencies, *Computers in Human Behavior: Artificial Humans*, Vol. 1, No. 2, 100014.
- Enholm, I.M., Papagiannidis, E., Mikalef, P. & Krogstie, J. (2022), Artificial Intelligence and Business Value: A Literature Review, *Information Systems Frontiers*, Vol. 24, No. 5, 1709-1734.
- Felzmann, H., Villaronga, E.F., Lutz, C. & Tamò-Larrieux, A. (2020), Towards Transparency by Design for Artificial Intelligence: H. Felzmann et al. , *Science and Engineering Ethics*, Vol. 26, No. 6, 3333-3361.
- Fügener, A., Grahl, J., Gupta, A. & Ketter, W. (2022), Cognitive Challenges in Human-Artificial Intelligence Collaboration: Investigating the Path Toward Productive Delegation, *Information Systems Research*, Vol. 33, No. 2, 678-696.
- Galbraith, J.R. (1974), Organization Design: An Information Processing View, *Interfaces*, Vol. 4, No. 3, 28-36.
- Gebru, T., Morgenstern, J., Vecchione, B., Vaughan, J.W., Wallach, H., Daumé III, H. & Crawford, K. (2021), Datasheets for Datasets, *Communications of the ACM*, Vol. 64, No. 12, 86-92.
- Glikson, E. & Woolley, A.W. (2020), Human Trust in Artificial Intelligence: Review of Empirical Research, *Academy of Management Annals*, Vol. 14, No. 2, 627-660.
- Hayes, A.F. (2009), Beyond Baron and Kenny: Statistical Mediation Analysis in the New Millennium, *Communication Monographs*, Vol. 76, No. 4, 408-420.
- Herschel, M., Diestelkämper, R. & Ben Lahmar, H. (2017), A Survey on Provenance: What for? What Form? What From?, *The VLDB Journal*, Vol. 26, No. 6, 881-906.

- Hornberger, M., Bewersdorff, A. & Nerdel, C. (2023), What Do University Students Know About Artificial Intelligence? Development and Validation of an AI Literacy Test, *Computers and Education: Artificial Intelligence*, Vol. 5, 100165.
- Ji, Z., Lee, N., Frieske, R., Yu, T., Su, D., Xu, Y., Ishii, E., Bang, Y.J., Madotto, A. & Fung, P. (2023), Survey of Hallucination in Natural Language Generation, *ACM Computing Surveys*, Vol. 55, No. 12, Article 248, 1-38.
- Jia, N., Luo, X., Fang, Z. & Liao, C. (2024), When and How Artificial Intelligence Augments Employee Creativity, *Academy of Management Journal*, Vol. 67, No. 1, 5-32.
- Jussupow, E., Spohrer, K., Heinzl, A. & Gawlitza, J. (2021), Augmenting Medical Diagnosis Decisions? An Investigation Into Physicians' Decision-Making Process with Artificial Intelligence, *Information Systems Research*, Vol. 32, No. 3, 713-735.
- Kaškevičius, L. (2025), Application of Artificial Intelligence in the Health Insurance Claims Management Process, *Journal of Management Changes in the Digital Era*, Vol. 2, No. 1, 148-159.
- Krakowski, S., Luger, J. & Raisch, S. (2023), Artificial Intelligence and the Changing Sources of Competitive Advantage, *Strategic Management Journal*, Vol. 44, No. 6, 1425-1452.
- Lakens, D. (2022), Sample Size Justification, *Collabra: Psychology*, Vol. 8, No. 1, 33267.
- Laupichler, M.C., Aster, A., Haverkamp, N. & Raupach, T. (2023), Development of the Scale for the Assessment of Non-Experts' AI Literacy: An Exploratory Factor Analysis, *Computers in Human Behavior Reports*, Vol. 12, 100338.
- Lebovitz, S., Lifshitz-Assaf, H. & Levina, N. (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*, Vol. 33, No. 1, 126-148.
- Lee, J.D. & See, K.A. (2004), Trust in Automation: Designing for Appropriate Reliance, *Human Factors: The Journal of the Human Factors and Ergonomics Society*, Vol. 46, No. 1, 50-80.
- Lintner, T. (2024), A Systematic Review of AI Literacy Scales, *npj Science of Learning*, Vol. 9, No. 1, 50.
- Mäntymäki, M., Minkkinen, M., Birkstedt, T. & Viljanen, M. (2022), Defining Organizational AI Governance, *AI and Ethics*, Vol. 2, No. 4, 603-609.
- Mikalef, P. & Gupta, M. (2021), Artificial Intelligence Capability: Conceptualization, Measurement Calibration, and Empirical Study on Its Impact on Organizational Creativity and Firm Performance, *Information & Management*, Vol. 58, No. 3, 103434.
- Mökander, J., Morley, J., Taddeo, M. & Floridi, L. (2021), Ethics-Based Auditing of Automated Decision-Making Systems: Nature, Scope, and Limitations: J. Mökander et al. , *Science and Engineering Ethics*, Vol. 27, No. 4, Article 44.
- Morris, T.P., White, I.R. & Crowther, M.J. (2019), Using Simulation Studies to Evaluate Statistical Methods, *Statistics in Medicine*, Vol. 38, No. 11, 2074-2102.
- Ng, D.T.K., Leung, J.K.L., Chu, S.K.W. & Qiao, M.S. (2021), Conceptualizing AI Literacy: An Exploratory Review, *Computers and Education: Artificial Intelligence*, Vol. 2, 100041.
- Noy, S. & Zhang, W. (2023), Experimental Evidence on the Productivity Effects of Generative Artificial Intelligence, *Science*, Vol. 381, No. 6654, 187-192.

Papagiannidis, E., Enholm, I.M., Dremel, C., Mikalef, P. & Krogstie, J. (2023), Toward AI Governance: Identifying Best Practices and Potential Barriers and Outcomes, *Information Systems Frontiers*, Vol. 25, No. 1, 123-141.

Papagiannidis, E., Mikalef, P. & Conboy, K. (2025), Responsible Artificial Intelligence Governance: A Review and Research Framework, *The Journal of Strategic Information Systems*, Vol. 34, No. 2, 101885.

Pinski, M. & Benlian, A. (2024), AI Literacy for Users: A Comprehensive Review and Future Research Directions of Learning Methods, Components, and Effects, *Computers in Human Behavior: Artificial Humans*, Vol. 2, No. 1, 100062.

Raisch, S. & Krakowski, S. (2021), Artificial Intelligence and Management: The Automation-Augmentation Paradox, *Academy of Management Review*, Vol. 46, No. 1, 192-210.

Shneiderman, B. (2020), Bridging the Gap Between Ethics and Practice: Guidelines for Reliable, Safe, and Trustworthy Human-Centered AI Systems, *ACM Transactions on Interactive Intelligent Systems*, Vol. 10, No. 4, Article 26, 1-31.

Vaccaro, M., Almaatouq, A. & Malone, T. (2024), When Combinations of Humans and AI Are Useful: A Systematic Review and Meta-Analysis, *Nature Human Behaviour*, Vol. 8, No. 12, 2293-2303.

Vasconcelos, H., Jörke, M., Grunde-McLaughlin, M., Gerstenberg, T., Bernstein, M.S. & Krishna, R. (2023), Explanations Can Reduce Overreliance on AI Systems During Decision-Making, *Proceedings of the ACM on Human-Computer Interaction*, Vol. 7, No. CSCW1, 1-38.

Walters, W.H. & Wilder, E.I. (2023), Fabrication and Errors in the Bibliographic Citations Generated by ChatGPT, *Scientific Reports*, Vol. 13, No. 1, 14045.

Wang, B., Rau, P.L.P. & Yuan, T. (2023), Measuring User Competence in Using Artificial Intelligence: Validity and Reliability of Artificial Intelligence Literacy Scale, *Behaviour & Information Technology*, Vol. 42, No. 9, 1324-1337.