

**RESPONSIBLE ARTIFICIAL INTELLIGENCE AND ORGANIZATIONAL
PERFORMANCE: THE MEDIATING ROLE OF AI GOVERNANCE
CAPABILITY AND THE MODERATING ROLE OF ETHICAL LEADERSHIP**

M. Saddique Khan

BS Student, National College of Business Administration and Economics, Lahore.

m.saddique2019@gmail.com

Muneeb Ullah

MS Scholar, Department of Business Administration, University of Punjab, Lahore.

Muneeb.msscholar@yahoo.com

ABSTRACT

The rapid diffusion of artificial intelligence (AI) into core organizational processes has outpaced firms' capacity to govern it responsibly, generating a widening gap between AI *adoption* and AI *accountability*. While an emerging stream of research links responsible AI (RAI) practices to firm-level outcomes, the mechanisms through which ethical AI principles translate into measurable organizational performance remain theoretically underspecified and empirically fragmented. Drawing on dynamic capabilities theory and upper echelons theory, this paper develops an integrative model in which **AI governance capability** — a firm's socio-technical capacity to sense AI-related risk, embed fairness, transparency, and accountability standards into AI systems, and reconfigure decision rights accordingly — mediates the relationship between responsible AI orientation and organizational performance. We further theorize that **ethical leadership** moderates this indirect effect, strengthening the translation of responsible AI principles into governance capability and organizational outcomes by legitimizing ethical discourse, reducing employee resistance, and aligning managerial incentives with responsible conduct. The paper synthesizes and critically extends recent empirical work (2021–2026) on responsible AI governance, AI-enabled dynamic capabilities, and ethical leadership in technology-intensive contexts, identifying a persistent gap: existing studies examine RAI's performance effects largely as direct or singly-mediated relationships, without accounting for the boundary condition that leadership ethicality imposes on governance institutionalization. We propose a moderated-mediation model, a multi-source, time-lagged survey design (targeting mid-to-large firms undergoing AI-enabled digital transformation), and a partial least squares structural equation modeling (PLS-SEM) analytical strategy suited to complex, formative-reflective constructs. We articulate directional hypotheses and discuss anticipated patterns of results, their theoretical interpretation against dynamic capabilities and stakeholder theory, and practical implications for AI governance design. The paper contributes a testable, boundary-condition-sensitive model that reframes responsible AI from a compliance cost to a performance-relevant dynamic capability, contingent on leadership ethicality.

Keywords: responsible AI, AI governance capability, ethical leadership, organizational performance, dynamic capabilities, moderated mediation

1. Introduction

1.1 Background and Context

Artificial intelligence has moved from a peripheral IT investment to a constitutive element of organizational strategy, decision-making, and value creation (Xiong et al., 2026; Fosso Wamba et al., 2024). As AI systems increasingly mediate hiring, credit, pricing, and operational decisions, organizations face mounting pressure — regulatory, reputational, and normative — to ensure that these systems are fair, transparent, accountable, and aligned with human welfare (Díaz-Rodríguez et al., 2023; Barredo Arrieta et al., 2020). This pressure has crystallized into the construct of **responsible AI (RAI)**: the deliberate integration of ethical principles into the design, deployment, and oversight of AI systems, intended to function as institutional "guardrails" against misuse and unintended harm (Xiong et al., 2026; Shneiderman, 2021).

Industry evidence suggests that RAI is no longer merely defensive. PricewaterhouseCoopers' 2025 global survey found that a majority of executives now attribute measurable returns — in efficiency, customer experience, and cybersecurity posture — to responsible AI investment (PwC, 2025). Yet the mechanisms by which RAI principles are converted into organizational performance remain poorly theorized. Most existing accounts either (a) treat RAI as a direct, unmediated predictor of performance, implicitly assuming that ethical intent alone is sufficient, or (b) study isolated ethical sub-dimensions — algorithmic bias, explainability, data privacy — without integrating them into a coherent organizational capability (Xia, Chen, Zhang, & Kamal, 2026; Kumar, Dwivedi, & Anand, 2023).

1.2 Research Gap

Three gaps motivate this study. First, the literature lacks a unifying **mediating mechanism** that explains *how* RAI principles are operationalized into performance-relevant organizational action. Recent work has begun to address this: Xia et al. (2026) demonstrate that responsible AI governance affects corporate performance in a quasi-natural experimental setting, and Papagiannidis, Mikalef, Krogstie, and Conboy (2022) show that knowledge management capabilities mediate the RAI governance–competitive performance link. However, these studies conceptualize the mediating capability narrowly (e.g., as knowledge management) rather than as an integrated *governance* capability encompassing risk-sensing, ethical embedding, and structural reconfiguration — the tripartite logic central to dynamic capabilities theory (Teece, 2007).

Second, the literature has not adequately theorized **boundary conditions**. Organizational behavior research shows that ethical leadership shapes how employees interpret and respond to AI-driven organizational change — for instance, ethical leadership moderates the relationship between AI adoption and employee psychological outcomes (Nature Humanities & Social Sciences Communications, 2025, on AI adoption, psychological safety, and depression). Yet this moderating logic has not been extended to the firm-level RAI–performance relationship, leaving unanswered whether the *same* level of stated RAI commitment yields different performance returns depending on the ethicality of the leadership enacting it.

Third, much of the extant empirical base is either single-country (e.g., Chinese listed manufacturing firms; Lebanese SMEs) or industry-specific (healthcare, public sector), limiting generalizability of proposed mechanisms across institutional contexts (Kumar et al., 2023; study of Lebanese SMEs in PMC, 2024).

1.3 Objectives and Research Questions

This paper aims to (1) theoretically integrate responsible AI, AI governance capability, and organizational performance into a coherent, testable moderated-mediation model; (2) specify ethical leadership as a boundary condition on the RAI–performance relationship; and (3) provide a rigorous empirical design capable of testing this model across firms undergoing AI-enabled transformation. Formally:

- **RQ1:** Does AI governance capability mediate the relationship between responsible AI orientation and organizational performance?
- **RQ2:** Does ethical leadership moderate the strength of the indirect (mediated) effect of responsible AI orientation on organizational performance via AI governance capability?
- **RQ3:** At what stage of the causal chain (RAI→capability, or capability→performance) does ethical leadership exert its strongest moderating influence?

1.4 Significance of the Study

Theoretically, the paper repositions responsible AI from a compliance-oriented cost center to a **performance-relevant dynamic capability**, while simultaneously introducing leadership ethicality as a contingency that determines whether this capability is realized. Practically, the model offers executives a diagnostic logic: RAI principles alone are insufficient: they must be embedded into governance structures, and this embedding is itself conditional on the ethical credibility of leadership. For policymakers and standard-setters, the model clarifies why uniform RAI regulatory mandates may produce heterogeneous performance outcomes across firms with differing leadership cultures.

2. Literature Review

2.1 Responsible AI and Organizational Performance

The RAI literature has matured rapidly since 2021. Early conceptual work (Barredo Arrieta et al., 2020) established fairness, transparency, accountability, and privacy as RAI's constitutive pillars. Subsequent empirical studies have tested RAI's performance implications with increasing sophistication. Xia et al. (2026), using a quasi-natural experimental design, provide causal-adjacent evidence that responsible AI governance improves corporate performance, while explicitly noting that "the multidimensional impact of responsible AI governance...on firm performance holistically and in an integrated manner" remains under-examined — a gap this paper directly targets. Xiong et al. (2026), drawing on the attention-based view, show that top-management-team attention to RAI shapes firm innovation outcomes, implying that RAI's performance effects are cognitively and structurally mediated rather than direct. Kumar et al. (2023), in a healthcare context, find that RAI's effect on market performance operates through patients' cognitive engagement — an early, domain-specific instance of the mediation logic this paper generalizes.

Critical assessment: these studies collectively establish that RAI matters for performance, but each mediates the relationship through a construct specific to its empirical context (patient engagement, innovation attention, knowledge management). No study has yet tested a *governance capability* mediator general enough to travel across industries while remaining theoretically grounded in dynamic capabilities logic.

2.2 AI Governance Capability as a Dynamic Capability

Dynamic capabilities theory (Teece, 2007; Teece, Peteraf, & Leih, 2016) conceptualizes firm capability as the capacity to sense opportunities and threats, seize them through resource reconfiguration, and transform organizational structures accordingly. Recent scholarship has begun applying this lens specifically to AI governance. Papagiannidis, Enholm, Dremel, Mikalef, and Krogstie (2022) identify best practices and barriers in AI governance institutionalization, framing governance as an emergent organizational capability rather than a static policy artifact. Building on this, Papagiannidis et al. (2022) demonstrate that RAI governance's effect on competitive performance is mediated by knowledge management capability — direct empirical precedent for the mediation architecture proposed here, though narrower in scope. Fosso Wamba, Queiroz, and Trinchera (2024) similarly show that AI-enabled dynamic capability affects environmental performance through a data-driven culture mediator, reinforcing that AI's performance effects are rarely direct but channeled through intermediate organizational capacities.

Critical assessment: the dynamic-capabilities-as-mediator logic is well established for *adjacent* constructs (knowledge management, data-driven culture) but has not been specified for *governance* capability per se — i.e., the capacity to translate ethical AI principles into enforceable risk-sensing routines, accountability structures, and decision-rights reallocation. This paper's first contribution is to isolate and operationalize this governance-specific capability as the mediating mechanism.

2.3 Ethical Leadership as a Boundary Condition

Ethical leadership — characterized by normatively appropriate conduct, two-way communication, and reinforcement of ethical standards (Brown, Treviño, & Harrison's foundational framework, as extended in recent AI-context studies) — has recently been examined as a moderator of AI's organizational effects. A 2025 three-wave study of South Korean employees found that ethical leadership moderates the relationship between AI adoption and employee psychological safety, buffering AI's adverse effects on wellbeing (Humanities and Social Sciences Communications, 2025). Complementary work on paradoxical leadership shows that leadership style conditions whether AI adoption translates into positive employee outcomes such as knowledge sharing (Hu, Gao, Agafari, Zhang, & Cong, 2025). At the strategic level, qualitative and practitioner-oriented literature (e.g., ISACA's 2024 analysis of RAI deployment; the World Economic Forum's 2024 responsible-AI playbook) consistently identifies leadership commitment as the primary bottleneck between RAI aspiration and RAI practice, though this observation has rarely been formalized as a statistical moderator in firm-level performance models.

Critical assessment: existing moderation evidence is concentrated at the *individual/psychological* level (employee wellbeing, knowledge sharing) rather than the *firm/strategic* level (governance capability, organizational performance). This paper's second contribution is to elevate ethical leadership from an individual-level buffer to a firm-level boundary condition on the RAI–governance capability–performance causal chain.

2.4 Synthesis and Positioning of the Present Study

Table 1 summarizes how this paper advances the reviewed literature.

Table 1. Positioning Relative to Recent Literature (2021–2026)

Study	Mediator/Moderator examined	Level of analysis	Gap addressed by present study
Xia et al. (2026)	None (direct RAI → governance performance)	Firm	Introduces governance-capability mediator
Xiong et al. (2026)	TMT attention (mediator)	Firm/TMT	Reframes mediator as an operational capability, not attention
Papagiannidis et al. (2022)	Knowledge management capability (mediator)	Firm	Generalizes mediator to broader governance capability construct
Kumar et al. (2023)	Patient cognitive engagement (mediator)	Individual/healthcare	Extends to cross-industry firm-level mechanism
HSSC (2025), Hu et al. (2025)	Ethical/paradoxical leadership (moderator)	Individual/employee	Elevates moderator to firm-level performance model

This positioning motivates the following conceptual model and hypotheses.

3. Conceptual Model and Hypotheses

Figure 1 (described): Responsible AI Orientation (independent variable) → AI Governance Capability (mediator) → Organizational Performance (dependent variable), with Ethical Leadership moderating both the RAI→Governance Capability path (a-path) and the Governance Capability→Performance path (b-path), consistent with a moderated-mediation (Model 59-type, Hayes PROCESS framework) design.

- **H1:** Responsible AI orientation is positively associated with AI governance capability.
- **H2:** AI governance capability is positively associated with organizational performance.
- **H3:** AI governance capability mediates the positive relationship between responsible AI orientation and organizational performance.
- **H4a:** Ethical leadership positively moderates the relationship between responsible AI orientation and AI governance capability, such that the relationship is stronger when ethical leadership is high.
- **H4b:** Ethical leadership positively moderates the relationship between AI governance capability and organizational performance, such that the relationship is stronger when ethical leadership is high.
- **H5:** The indirect effect of responsible AI orientation on organizational performance via AI governance capability is conditional on ethical leadership (moderated mediation), such that the indirect effect is stronger at higher levels of ethical leadership.

4. Methodology

4.1 Research Design

We propose a **quantitative, cross-sectional-supplemented-by-time-lagged, multi-source survey design**, appropriate for testing moderated-mediation relationships among latent organizational constructs (cf. the 417-respondent SME design in PMC, 2024, and the 381-respondent, 3-wave design in HSSC, 2025, both of which validate this design family for AI-organizational-behavior research). A single quantitative design is justified over a mixed-methods approach because the model's core contribution is the *statistical* test of mediation and moderated mediation; a supplementary qualitative phase (semi-structured interviews with Chief AI/Data Officers) is proposed as an optional Phase 2 for theoretical elaboration but is not required to test H1–H5.

4.2 Sample and Data Collection

- **Target population:** Mid-to-large organizations (≥ 250 employees) across finance, healthcare, manufacturing, and professional services that have formally adopted AI in at least one core business function, sampled across at least two national institutional contexts to address the single-country limitation identified in Section 2.
- **Respondents (multi-source, to reduce common-method bias):** (1) senior AI/data governance leads or CIOs, reporting on RAI orientation and governance capability; (2) direct-report middle managers, reporting on perceived ethical leadership; (3) archival/financial data (ROA, revenue growth, or Balanced Scorecard indices) for organizational performance, reducing reliance on self-reported performance.
- **Target sample size:** A minimum of 300–400 firm-level responses (following power analysis for PLS-SEM models with moderated mediation, using G*Power or the inverse square root method of Kock and Hadaya, 2018), consistent with comparable published designs (e.g., $n = 417$ in PMC, 2024; $n = 381$ in HSSC, 2025).
- **Time structure:** A two-wave, three-month-lagged design (RAI orientation and governance capability at T1; performance outcomes at T2) to mitigate reverse-causality concerns.

4.3 Measurement

- **Responsible AI orientation:** adapted multi-item reflective scale drawing on fairness, transparency, accountability, and privacy sub-dimensions (Barredo Arrieta et al., 2020; Díaz-Rodríguez et al., 2023).
- **AI governance capability:** newly developed formative-reflective (Type II) construct capturing risk-sensing routines, ethical-principle embedding into system design, and decision-rights reconfiguration, informed by Papagiannidis et al.'s (2022) AI governance best-practice taxonomy and Teece's (2007) sensing-seizing-transforming logic.
- **Ethical leadership:** established scale (10-item Ethical Leadership Scale, Brown et al. tradition), rated by subordinate managers.
- **Organizational performance:** composite of financial (ROA, revenue growth) and non-financial (innovation output, stakeholder trust index) indicators, following calls for multidimensional performance measurement in the RAI literature (Xia et al., 2026).
- **Controls:** firm size, industry, AI maturity stage, national regulatory stringency (e.g., presence of an AI Act-equivalent regime).

4.4 Analytical Strategy

We propose **partial least squares structural equation modeling (PLS-SEM)**, justified by (a) the presence of formative indicators in the AI governance capability construct, (b)

Pakistan Journal of Management & Social Science

<https://pakistanjournalofmanagement.com/index.php/Journal>

a theory-development (rather than theory-confirmation) research objective, and (c) precedent in closely related studies (e.g., HSSC, 2025, uses structural equation modeling for a directly analogous AI-adoption/ethical-leadership model). Moderated mediation (H5) would be tested using Hayes' PROCESS macro (Model 59 logic, adapted for SEM) or the two-stage approach recommended by Memon et al. (2019) for PLS-SEM moderated-mediation, with bootstrapped confidence intervals (10,000 resamples) for indirect-effect significance testing. Common method bias would be assessed via Harman's single-factor test and a marker-variable technique, mitigated structurally through the multi-source data collection described above.

4.5 Justification of Choices

The multi-source, time-lagged design directly addresses two recurrent limitations flagged in the reviewed literature: single-respondent common-method bias and cross-sectional causal ambiguity. The choice of PLS-SEM over covariance-based SEM reflects the exploratory, capability-construct-development nature of this study, consistent with methodological guidance for emergent organizational capability research (Hair, Risher, Sarstedt, & Ringle, 2019).

5. Hypothesized Results and Analytical Expectations

(Presented as theoretically derived expectations to guide empirical testing, not as findings from a conducted study.)

Based on the theoretical model and the directional logic of the reviewed literature, we would expect the structural model to indicate:

1. A significant positive path from responsible AI orientation to AI governance capability (supporting H1), consistent with the logic that stated ethical commitment is a necessary precondition for governance institutionalization (Papagiannidis et al., 2022).
2. A significant positive path from AI governance capability to organizational performance (supporting H2), consistent with Xia et al. (2026) and Fosso Wamba et al. (2024).
3. A significant indirect effect (RAI orientation $\rightarrow$ governance capability $\rightarrow$ performance) with a bootstrapped confidence interval excluding zero, indicating full or partial mediation (supporting H3). Given that RAI orientation may also retain a modest direct effect on performance (e.g., via reputational signaling to markets), **partial** rather than full mediation is the theoretically more plausible pattern.
4. A significant, positive interaction term (RAI orientation $\times$ Ethical Leadership) predicting governance capability, indicating that the a-path strengthens as ethical leadership rises (supporting H4a) — consistent with the buffering/amplifying role ethical leadership plays in analogous individual-level AI studies (HSSC, 2025).
5. A plausible but theoretically weaker interaction effect on the b-path (governance capability $\times$ Ethical Leadership $\rightarrow$ performance) (H4b), since once governance capability is institutionalized, its performance effects may be less contingent on leadership style than its initial formation was — an asymmetry worth explicit empirical attention rather than an assumed symmetric moderation.
6. An index-of-moderated-mediation with a confidence interval excluding zero, indicating that the strength of the indirect effect varies significantly across levels of ethical leadership (supporting H5), with the conditional indirect effect strongest at +1 SD ethical leadership and attenuated (though not necessarily null) at -1 SD.

We would recommend visualizing the conditional indirect effects via a Johnson-Neyman floodlight analysis to identify the precise threshold of ethical leadership at

which the RAI→performance indirect effect becomes statistically significant, which is likely to be of greater practical value to managers than a simple high/low split.

6. Discussion

6.1 Interpretation

If the hypothesized pattern is empirically confirmed, the model would suggest that responsible AI's performance benefits are neither automatic nor uniformly distributed: they are *manufactured* through governance capability, and that manufacturing process is itself accelerated or throttled by the ethical credibility of leadership. This reframes a normative construct (responsible AI) as a strategic resource whose value is capability-mediated and leadership-contingent — extending dynamic capabilities theory into the ethics-of-AI domain, and extending upper echelons theory (Hambrick & Mason, 1984; operationalized recently via the attention-based view in Xiong et al., 2026) by specifying leadership *ethicality*, rather than mere attention or cognition, as a performance-relevant executive characteristic.

6.2 Comparison with Existing Literature

The expected partial-mediation pattern would align with, and generalize, Papagiannidis et al.'s (2022) knowledge-management-capability mediation finding, suggesting that governance capability may be a superordinate construct of which knowledge management is one component. The expected asymmetric moderation (stronger at the RAI→capability stage than the capability→performance stage) would extend, rather than merely replicate, the individual-level ethical leadership moderation found in HSSC (2025): whereas that study shows ethical leadership buffers *employees* against AI's negative effects, our model would show ethical leadership *accelerates organizational learning and institutionalization* — a structurally different mechanism operating at a different level of analysis.

6.3 Theoretical Implications

The paper contributes a boundary-condition-sensitive extension of dynamic capabilities theory to the RAI domain, addressing calls (Xia et al., 2026) for integrated, multidimensional treatments of RAI governance. It also contributes to leadership theory by demonstrating (pending empirical confirmation) that ethical leadership's organizational relevance extends beyond employee wellbeing outcomes into capability formation and firm performance — a multi-level bridge rarely constructed explicitly in the AI-management literature.

6.4 Practical Implications

For executives, the model implies that RAI investment without parallel investment in ethical leadership development is likely to yield attenuated returns. For boards and governance committees, it suggests that AI governance capability should be treated as a measurable organizational asset — auditable, benchmarkable, and reportable — rather than a compliance checklist. For policymakers, it implies that regulatory mandates for RAI disclosure (e.g., EU AI Act-adjacent regimes) may generate heterogeneous firm-level performance effects depending on unobserved leadership quality, a consideration relevant to impact assessment of such regulation.

7. Conclusion

7.1 Summary of Key Contributions

This paper develops and operationalizes a moderated-mediation model linking responsible AI orientation to organizational performance through AI governance capability, contingent on ethical leadership. It synthesizes and critically extends recent

Pakistan Journal of Management & Social Science

<https://pakistanjournalofmanagement.com/index.php/Journal>

(2021–2026) empirical literature, identifies a specific theoretical gap — the absence of an integrated, leadership-contingent governance-capability mediator — and proposes a rigorous, previously unimplemented empirical design capable of testing it.

7.2 Limitations

As a conceptual paper with a proposed (not yet executed) empirical design, the central limitation is the absence of primary data; the hypothesized results in Section 5 are theoretically derived expectations, not empirical findings, and must be treated accordingly until tested. The proposed governance-capability scale requires formal validation (exploratory and confirmatory factor analysis) prior to use. Cross-national data collection, while strengthening generalizability, introduces measurement-invariance challenges that would need to be formally tested (e.g., via multi-group SEM invariance testing) before pooling data.

7.3 Future Research Directions

Future work should (1) empirically field the proposed design; (2) explore additional boundary conditions, such as national AI regulatory stringency or organizational AI maturity stage, as potential three-way moderators; (3) examine reverse or reciprocal dynamics, whereby strong organizational performance enables further investment in RAI and governance capability, using longitudinal panel designs; and (4) extend the model to agentic and autonomous AI systems, where governance capability requirements are likely to intensify further (PwC, 2025).

References

- Ahmed, M., Almotairi, M. A., Ullah, S., & Alam, A. (2014). Mobile banking adoption: a qualitative approach towards the assessment of TAM model in an emerging economy. *Academic Research International*, 5(6), 248.
- Alam, A., Malik, O. M., Ahmed, M., & Gaadar, K. (2015). Empirical analysis of tourism as a tool to increase foreign direct investment in developing country: Evidence from Malaysia. *Mediterranean Journal of Social Sciences*, 6(4), 201–206.
- Barredo Arrieta, A., Díaz-Rodríguez, N., Del Ser, J., Bennetot, A., Tabik, S., Barbado, A., ... & Herrera, F. (2020). Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information Fusion*, 58, 82–115.
- Díaz-Rodríguez, N., Del Ser, J., Coeckelbergh, M., López de Prado, M., Herrera-Viedma, E., & Herrera, F. (2023). Connecting the dots in trustworthy Artificial Intelligence: From AI principles, ethics, and key requirements to responsible AI systems and regulation. *Information Fusion*, 99, 101896.
- Fosso Wamba, S., Queiroz, M. M., & Trinchera, L. (2024). The role of artificial intelligence-enabled dynamic capability on environmental performance: The mediation effect of a data-driven culture in France and the USA. *International Journal of Production Economics*, 268, 109131.
- Hair, J. F., Risher, J. J., Sarstedt, M., & Ringle, C. M. (2019). When to use and how to report the results of PLS-SEM. *European Business Review*, 31(1), 2–24.
- Hidayat-ur-Rehman, I., Ahmad, A., Ahmed, M., & Alam, A. (2021). Mobile Applications to Fight against COVID-19 Pandemic: The Case of Saudi Arabia. *TEM Journal*, 10(1).

Pakistan Journal of Management & Social Science

<https://pakistanjournalofmanagement.com/index.php/Journal>

- Hu, X., Gao, H., Agafari, T., Zhang, M. Q., & Cong, R. (2025). How and when artificial intelligence adoption promotes employee knowledge sharing? The role of paradoxical leadership and technophilia. *Frontiers in Psychology*, 16, 1573587.
- Humanities and Social Sciences Communications. (2025). The dark side of artificial intelligence adoption: Linking artificial intelligence adoption to employee depression via psychological safety and ethical leadership. *Humanities and Social Sciences Communications*, 12, article 05040.
- Kock, N., & Hadaya, P. (2018). Minimum sample size estimation in PLS-SEM: The inverse square root and gamma-exponential methods. *Information Systems Journal*, 28(1), 227–261.
- Kumar, P., Dwivedi, Y. K., & Anand, A. (2023). Responsible artificial intelligence (AI) for value formation and market performance in healthcare: The mediating role of patient's cognitive engagement. *Information Systems Frontiers*, 25(6), 2197–2220.
- Memon, M. A., Cheah, J. H., Ramayah, T., Ting, H., Chuah, F., & Cham, T. H. (2019). Moderation analysis: Issues and guidelines. *Journal of Applied Structural Equation Modeling*, 3(1), i–xi.
- Papagiannidis, E., Enholm, I. M., Dremel, C., Mikalef, P., & Krogstie, J. (2022). Toward AI governance: Identifying best practices and potential barriers and outcomes. *Information Systems Frontiers*, 25, 1–19.
- Papagiannidis, E., Mikalef, P., Krogstie, J., & Conboy, K. (2022). From responsible AI governance to competitive performance: The mediating role of knowledge management capabilities. In *Responsible AI and Analytics for an Ethical and Inclusive Digitized Society* (pp. 1–13). Springer.
- PricewaterhouseCoopers. (2025). *PwC's 2025 Responsible AI survey: From policy to practice*. PwC.
- Teece, D. J. (2007). Explicating dynamic capabilities: The nature and microfoundations of (sustainable) enterprise performance. *Strategic Management Journal*, 28(13), 1319–1350.
- Teece, D., Peteraf, M., & Leih, S. (2016). Dynamic capabilities and organizational agility: Risk, uncertainty, and strategy in the innovation economy. *California Management Review*, 58(4), 13–35.
- Xia, H., Chen, H., Zhang, J. Z., & Kamal, M. M. (2026). Exploring the impact of responsible AI governance on corporate performance: A quasi-natural experiment. *Technological Forecasting and Social Change*, 223, 124425.
- Xiong, [initials], et al. (2026). Responsible artificial intelligence attention and firm innovation: An attention-based view. *Journal of Product Innovation Management*.