

The Influence of Reliable Model Operations on Trust and The Adoption of AI Systems

¹Sateesh Gottumukkala, ²S Shyam Prasad

¹Corresponding Author: Research Scholar, International School of Management Excellence (ISME), Bengaluru, INDIA.

²Professor, International School of Management Excellence (ISME), Bengaluru, INDIA.

Abstract

As Artificial Intelligence (AI) systems become increasingly integrated into business operations and daily life, the challenge of building trustworthy AI has emerged as a critical determinant of successful adoption. While significant attention was given to algorithmic innovation, the operational frameworks that govern AI model lifecycle management, collectively known as Model Operations (ModelOps), remain an underexplored factor in shaping stakeholder trust in AI systems. The study of these operational frameworks becomes even more critical in the era of generative and agentic AI. This study investigates how various ModelOps practices shape trust in AI systems through a quantitative survey of 266 AI professionals, business stakeholders and end users across diverse industries. Using validated measurement instruments, we examine how different dimensions of ModelOps influence trust in AI systems. Our findings reveal that organisations with reliable model operations are associated with higher levels of stakeholder trust in AI systems ($\beta = 0.46$, $p < 0.002$). In turn, trust significantly predicts both perceived value ($\beta = 0.44$, $p < 0.001$) and actual AI usage ($\beta = 0.91$, $p < 0.001$). Notably, trust was found to fully mediate the relationship between operational maturity and adoption outcomes, suggesting that technical reliability alone is insufficient to drive usage without first establishing a foundation of trust. These findings establish that in the era of generative and agentic AI, investing in robust ModelOps practices, specifically monitoring, quality, and human-in-the-loop oversight, is as crucial as algorithmic refinement for realising sustained organisational impact.

Keywords: ModelOps, MLOps, AI trust, Artificial Intelligence Adoption, Machine Learning Operations, trustworthy AI, Model Governance

1. Introduction

With the evolution of generative and agentic capabilities, artificial intelligence (AI) has transitioned from a futuristic concept to a present-day operational reality, fundamentally reshaping industries and business processes worldwide. AI has outperformed conventional solutions in many business areas, such as Manufacturing, Healthcare, Transportation, Banking, and Retail, which has helped increase the use of AI methods in these areas (Bharati et al., 2024; Lingao, 2024). AI is also transforming functional areas like Marketing, Finance, Operations, Human Resources management, etc (Haleem et al., 2022; Lingao, 2024; Routray, 2024). While AI is transforming business and functional areas, the adoption of AI systems is still low in many organisations. Trust is foundational for the adoption and effective use of AI systems, as it influences user acceptance, satisfaction, and overall interaction quality with these technologies (Stanley & Dorton, 2023; Yang & Wibowo, 2022).

Trust consistently emerged as a fundamental prerequisite for user adoption and sustained engagement (Davis et al., 2022). In the context of AI, this relationship becomes even more critical due to the inherent complexity and opacity of many AI systems. Stakeholders, from internal users to external customers and regulators, will not embrace or rely on AI systems they do not trust. Without trust, the immense value of AI investments

remains unrealised (Thompson & Williams, 2023). The complexity of trust in AI is influenced by various factors, including technical, regulatory, social and organisational factors (Mehrotra et al., 2023; Vibhuti Choubisa & Divyansh Choubisa, 2024).

As organisations accelerate their AI adoption, moving from isolated experiments to enterprise-wide integration, the focus is shifting from the mere creation of algorithms to the complex challenge of their sustained and reliable operations. Beyond the algorithms themselves, the systematic processes governing the entire lifecycle of AI models, from deployment and monitoring to governance and risk management, play a hidden yet pivotal role in shaping perceptions of AI trustworthiness and, consequently, in driving its successful adoption (Martinez et al., 2024). ModelOps offers a broader, enterprise-level governance framework essential for managing the complexities of a diverse and growing portfolio of AI and decision models (IBM Research, 2023).

1.1 Research Objectives and Questions

The primary objective of this research is to empirically investigate the relationships between ModelOps and trust in AI systems. Specifically, this study seeks to answer the following research questions:

RQ1: To what extent do the reliable ModelOps practices predict stakeholder trust in its AI systems?

RQ2: How does trust influence the use of AI systems within organisations and influence perceived value in AI systems?

1.2 Theoretical Contributions and Practical Implications

This research makes several important theoretical contributions to the emerging literature on AI Model Operations and trust. First, it provides empirical validation of the conceptual relationships between operational practices and trust outcomes in the context of AI. Second, it advances our understanding of trust as a mediating mechanism in technology adoption, extending established theories to the AI domain.

From a practical perspective, this research provides insights for organisations seeking to improve their AI adoption outcomes. By identifying which operational practices influence trust and use of AI, the findings can guide resource allocation and strategic decision-making in AI initiatives.

2. Background and Related Work

The discipline of Model Operations, shortly known as ModelOps, has evolved rapidly in response to the growing complexity and scale of AI deployments in enterprise environments. Gartner defines ModelOps (or AI model operationalisation) as governance and life-cycle management of a wide range of artificial intelligence (AI) and decision models, including machine learning, knowledge graphs, rules, optimisation, linguistic and agent-based models. ModelOps has emerged in recent years as a unifying construct for describing the capabilities required to operationalise AI models reliably and at scale. Rather than referring solely to deployment tooling, contemporary accounts frame ModelOps as a multidimensional organisational capability that integrates technical infrastructure, continuous oversight, governance mechanisms and quality assurance across the model lifecycle (Amershi et al., 2019; Kreuzberger et al., 2023).

Synthesising insights from the literature, four core dimensions, such as technical infrastructure, operational vigilance, governance and compliance and output integrity, are the foundation for reliable ModelOps. First, the technical infrastructure dimension concerns the foundational pipelines and platforms that support production-grade model development and operations. Empirical case studies from large software and cloud providers highlight the importance of standardised tools, rigorous versioning, and automated deployment processes for managing complex machine-learning systems (Amershi et al., 2019; Breck et al., 2017). Sculley (Sculley et al., 2015) argued that ad hoc scripts and undocumented dependencies create “hidden technical

debt,” undermining reliability and making it difficult to reproduce or audit model behaviour. Collectively, this stream of work suggests that robust tooling, comprehensive version control and automated, repeatable deployment pipelines are necessary preconditions for dependable ModelOps.

Second, operational vigilance focuses on the ongoing oversight of models once deployed. A large body of work on concept and data drift demonstrates that model performance can deteriorate rapidly when underlying data distributions shift, especially in non-stationary environments (Gama et al., 2014; Lu et al., 2019). Production-oriented checklists, such as the ML Test Score (Breck et al., 2017) explicitly recommend continuous monitoring of input distributions, prediction quality, and service health to detect anomalies and performance decay. Together, these studies converge on the view that continuous monitoring and systematic drift detection are integral components of mature ModelOps, enabling timely intervention and adaptation as real-world conditions evolve.

Third, the governance and compliance dimension captures the structures that link technical operations to regulatory, ethical and organisational accountability requirements. AI literature emphasises auditability, traceability and human oversight as central to trustworthy AI development (Morley et al., 2020). Auditability through comprehensive logging, documentation and lineage tracking enables organisations to reconstruct which model was deployed, on which data and under which configuration, a capability increasingly demanded by emerging regulatory regimes. In parallel, human-in-the-loop (HIL) or “human-in-command” arrangements are widely promoted in ethical AI guidelines (Billeter et al., 2024) and human-centred AI (Shneiderman, 2020). These works argue that structured human review, approval gates, and override mechanisms can mitigate automation bias, support accountability and enhance stakeholder trust. This literature, therefore, motivates the inclusion of auditability and HIL approval as core pillars of ModelOps governance.

Finally, output integrity refers to the assurance that model predictions remain accurate, robust and fit for their intended purpose over time. Production-readiness rubrics for AI/ML systems foreground data validation, rigorous testing, and fairness and robustness checks as necessary safeguards against unexpected failure modes (Breck et al., 2017). At the same time, critical examinations of data and model bias demonstrate how seemingly minor quality issues in training data can lead to systematic harms and erode trust in AI systems (Suresh et al., 2020). This stream of research suggests that quality in ModelOps extends beyond point-in-time accuracy metrics to encompass robustness to distributional shift, resilience to noise and adherence to ethical and safety requirements. High standards of data and model quality are thus framed as essential for sustaining reliable and actionable predictions in operational settings.

Trust in AI systems has emerged as a critical research area, drawing from multiple disciplines including psychology, computer science, and organisational behaviour. The complexity of AI trust stems from its multifaceted nature, encompassing both cognitive and affective components that influence user behaviour and system adoption. Early work on trust and technology mainly emphasised model-level properties such as accuracy, usefulness, and transparency (Shin, 2021). However, emerging research argues that the operational environment in which models are engineered, deployed and maintained is equally critical for sustaining trust over time (Amershi et al., 2019; Billeter et al., 2024; Breck et al., 2017).

From an engineering perspective, studies show that brittle pipelines, informal tooling and weak configuration management lead to unpredictable behaviour in production, including silent failures and inconsistent results across environments (Breck et al., 2017; Sculley et al., 2015). Such instability erodes stakeholders’ willingness to rely on AI services, particularly in high-stakes settings. In contrast, mature ModelOps practices standardised tools, rigorous versioning of data and models, automated and test-driven deployments, enable reproducibility and controlled change management (Amershi et al., 2019; Zaharia et al., 2018). By reducing operational uncertainty and making behaviour more predictable, these practices support perceptions of reliability and competence, the two core dimensions of trust.

Operational vigilance further reinforces this relationship. The literature demonstrates that model performance

can degrade substantially when data distributions shift, even if the underlying algorithms remain unchanged (Gama et al., 2014; Lu et al., 2019). Continuous monitoring of input distributions, performance metrics, and fairness indicators enables timely detection and mitigation of such degradation (Breck et al., 2017; Kreuzberger et al., 2023). Responsible AI frameworks contend that these mechanisms are prerequisites for trustworthy deployment, because they limit the risk of undetected harms and signal that the organisation actively manages risks (Billeter et al., 2024). Governance and AI ethics consistently identify auditability, traceability, and human-in-the-loop (HIL) oversight as enabling trustworthy AI (Morley et al., 2020).

ModelOps acts as an enabler of trustworthy AI precisely because it operationalises principles such as accountability and transparency through concrete mechanisms, which include comprehensive logging, model lineage tracking, formal approval workflows, and documented human-in-loop arrangements. Human-centred AI research similarly argues that structured HIL controls and override capabilities reduce automation bias, align system behaviour with domain norms and thereby strengthen user and organisational trust (Shneiderman, 2020). Production-readiness rubrics and safety engineering approaches call for systematic validation of data quality, robustness, and fairness, in addition to conventional accuracy metrics (Breck et al., 2017; Varshney & Alemzadeh, 2017).

Taken together, the literature indicates that reliable ModelOps does not merely support efficient AI delivery but also shapes the perceived trustworthiness of AI systems by enhancing reproducibility, stability, oversight and accountability. Conceptually, ModelOps can therefore be viewed as a process-level antecedent of trust. It provides the organisational routines and safeguards that allow stakeholders to form confident expectations about how AI systems will behave in practice.

2.1 Theoretical Framework and Hypotheses

Based on the literature review, this research proposes a theoretical framework that positions Reliable ModelOps as a key antecedent to AI trust, which in turn influences adoption success. This framework draws from technology acceptance theory, trust theory, and operations management literature to explain the mechanisms through which operational practices influence adoption outcomes.

The framework suggests that ModelOps practices influence trust through multiple pathways. Reliability is enhanced through continuous monitoring and drift management practices that ensure consistent performance. Explainability is supported by comprehensive documentation and auditability practices that provide transparency into model behaviour. Fairness is promoted through governance practices that enforce bias detection and mitigation procedures. Safety is ensured through risk management practices that identify and mitigate potential harms.

Based on this theoretical framework, the following hypotheses are proposed:

Hypothesis 1 (H1): Reliable ModelOps practices (comprising tools, versioning, deployment, monitoring, drift management, auditability, HIL approval, and quality) have a significant positive effect on stakeholder Trust in AI systems.

Hypothesis 2 (H2): Stakeholder Trust in AI systems has a significant positive effect on the Perceived Value of those systems.

Hypothesis 3 (H3): Stakeholder Trust in AI systems has a significant positive effect on actual AI usage.

Hypothesis 4a (H4a): Trust significantly mediates the relationship between Reliable ModelOps and Perceived Value.

Hypothesis 4b (H4b): Trust significantly mediates the relationship between Reliable ModelOps and AI Usage.

3. Research Design & Methodology

3.1 Research Design

This study utilises a deductive, quantitative research design to investigate the structural linkages between operational maturity (ModelOps), psychological trust, and adoption outcomes. Given the complexity of the proposed model, which conceptualises Reliable ModelOps as a high-order latent construct, Structural Equation Modeling (SEM) was selected as the primary analytical engine. Specifically, we employed SPSS AMOS (v. 31) to conduct a covariance-based SEM. This choice was driven by the need for a confirmatory approach to validate the theoretically grounded paths between operational reliability and stakeholder trust, allowing for the simultaneous estimation of measurement error and structural coefficients

3.2 Sampling Strategy and Participant Profile

The empirical data were derived from a targeted survey of $N = 266$ AI and Machine Learning professionals. To reach this specialised cohort, a purposive non-probability sampling strategy was employed through professional networks and specialised technical communities. The final sample size satisfies the robust requirements for maximum likelihood estimation in AMOS, ensuring stable parameter estimation for a model of this complexity (Hair et al., 2019).

The respondent profile reflects a mature and technically proficient cohort. Age distribution was centred in the 31–40 (38%) and 21–30 (30.5%) brackets, with a significant portion of the sample possessing substantial industry tenure. 51.8% of respondents reported more than four years of direct AI experience, including 15.8% with over a decade of expertise. Professionally, the sample was dominated by Developers (35.7%) and those in Leadership roles (20.3%), followed by Foundation specialists (14.7%) and AI Engineers (13.2%). Geographically, while the sample included representation from North America (12.8%), Asia Pacific (12%), and Europe (10.9%), a vast majority of respondents (62.8%) operated within a global organisational context. Furthermore, the data reflects a heavy enterprise bias, with 65.4% of participants belonging to large-scale organisations with 10,000+ employees.

The target population for this study includes AI practitioners, business stakeholders, and end-users of AI systems across various industries and organisational contexts. To ensure adequate representation and statistical power, a target sample size of 266 respondents was established based on power analysis calculations for structural equation modeling with medium effect sizes (Cohen, 1988).

3.3 Survey Instrument Development

The survey instrument was developed through a systematic process involving literature review, expert consultation, and pilot testing. The final instrument consists of five main sections designed to capture the key constructs of interest while minimising respondent burden. Data was collected via a structured questionnaire comprising multi-item scales adapted from prior literature and industry frameworks. All items were measured using a five-point Likert scale ranging from “strongly disagree” to “strongly agree.” The key construct, Reliable ModelOps, was modeled as a second-order latent construct reflected by eight dimensions: Tools, Versioning, Deployment, Monitoring, Drift, Auditability, Human-in-the-Loop (HIL) Approval, and Quality. These dimensions were derived from established MLOps and AI governance literature (Billeter et al., 2024; Breck et al., 2017; Kreuzberger et al., 2023). The constructs of Trust, Perceived Value, and AI Usage were measured using validated scales adapted from prior studies on organisational trust and technology adoption (Davis, 1989; Mayer et al., 1995; Shin, 2021).

To ensure content validity and clarity, the questionnaire was pilot-tested with domain experts and minor refinements were made to improve item wording. Screening questions were included to confirm respondents' involvement in AI-related work, and attention checks were incorporated to ensure response quality.

3.4 Data Analysis

The data were analysed using Structural Equation Modeling (SEM) in AMOS to examine both the measurement and structural components of the proposed model. A two-step approach, consistent with Anderson and Gerbing (1988), was followed. In the first stage, Confirmatory Factor Analysis (CFA) was conducted to evaluate the measurement model. The higher-order construct of Reliable ModelOps was assessed through its observed dimensions. Convergent validity was examined through standardised factor loadings, with all indicators demonstrating acceptable and statistically significant loadings. Reliability and validity were primarily assessed through the strength and consistency of these loadings, given the SEM-based estimation framework. Model fit was evaluated using multiple goodness-of-fit indices, including the chi-square statistic (χ^2), the chi-square to degrees of freedom ratio (χ^2/df), Comparative Fit Index (CFI), and Root Mean Square Error of Approximation (RMSEA), following commonly accepted thresholds.

In the second stage, the structural model was estimated to test the hypothesised relationships among Reliable ModelOps, Trust, Perceived Value, and AI Usage. The model was estimated using Full Information Maximum Likelihood (FIML) to handle missing data, allowing the use of all available observations without requiring data imputation. Hypotheses were evaluated based on standardised path coefficients, critical ratios (C.R.) and associated significance levels.

The mediating role of Trust was assessed using the joint significance approach, wherein mediation is supported when both the path from the independent variable to the mediator and the path from the mediator to the dependent variable are statistically significant. Indirect effects were computed using the product of standardised coefficients to provide an estimate of mediation strength. This analytical strategy provides a rigorous and theory-driven evaluation of the proposed model, enabling the examination of how ModelOps maturity influences trust and, in turn, shapes AI adoption outcomes.

4. Results & Analysis

The structural equation model (SEM) was evaluated using Full Information Maximum Likelihood (FIML). This approach was selected to maintain the power of the analysis and avoid the biases inherent in listwise deletion or simple mean imputation, particularly given the complex, multi-indicator nature of the latent constructs. While the global fit indices ($\chi^2[34] = 265.484, p < .001$, CFI = .867, RMSEA = .160) suggest that the model represents a modest approximation of the observed covariance structure, the structural paths and factor loadings provide a robust foundation for examining the underlying theoretical relationships.

4.1 Measurement Model Results

The latent construct Reliable ModelOps was evaluated through eight standardised indicators. All factor loadings were statistically significant ($p < 0.001$) and ranged from .719 to .887. Specifically, Monitoring ($\beta = .887$), Quality ($\beta = .847$), and Human-in-the-loop Approval ($\beta = .832$) emerged as the strongest manifestations of operational reliability. The standardised loadings for Deployment ($\beta = .799$), Drift Management ($\beta = .799$), DevTools ($\beta = .744$), Auditability ($\beta = .723$), and Versioning ($\beta = .719$) further confirmed the convergent validity of the construct, suggesting that stakeholders associate ModelOps reliability with a holistic suite of governance and technical disciplines.

4.2 Structural Model & Hypothesis Testing

Following the validation of the measurement model, Path analysis was conducted to evaluate the relationships between ModelOps maturity, trust, and adoption. Hypothesis H1 was supported, as Reliable ModelOps exerted a significant positive effect on Trust ($\beta = .457, p = .002$). Trust, in turn, proved to be the most powerful antecedent of behavioural outcomes. Supporting H3, the relationship between Trust and AI Usage was exceptionally strong and significant ($\beta = .912, p < .001$). Furthermore, the path from Trust to

Perceived Value was positive and moderate ($\beta = .440$), supporting H2.

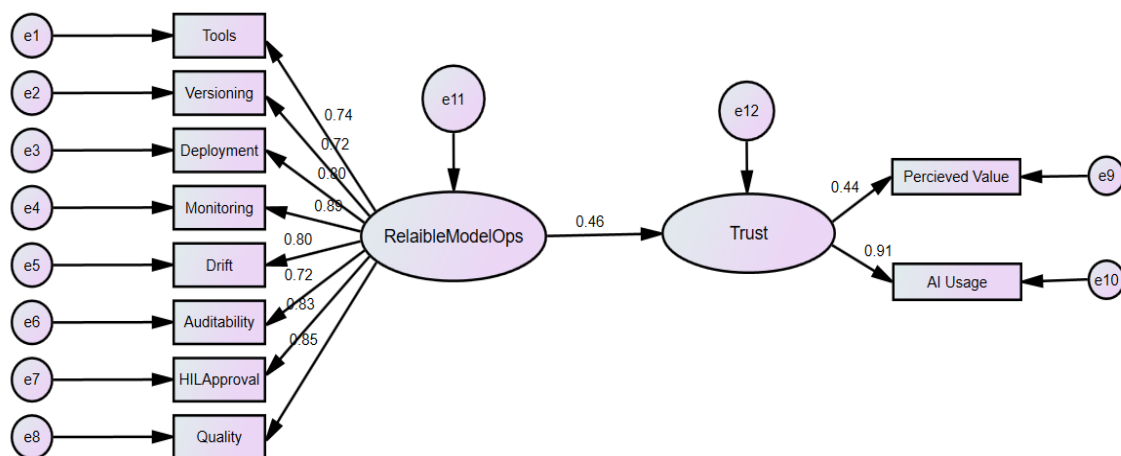


Figure 1: Structural Equation Model results showing standardised path coefficients.

Mediation was assessed using the joint significance method. Because the paths from Reliable ModelOps to Trust (Path a) and from Trust to the adoption outcomes (Path b) were both statistically significant, and mediation is established. The total indirect effect of ModelOps on AI Usage through Trust was substantial ($\beta = .417$), while the indirect effect on Perceived Value was $\beta = .201$. The absence of specified direct paths from ModelOps to the outcomes, coupled with the joint significance of these indirect routes, suggests a model of Full Mediation.

5. Discussion

The results of this study offer a compelling empirical bridge between the technical rigours of machine learning engineering and the behavioural domain of organisational adoption.

5.1 The Operational Foundations of Trust

A defining contribution of this research is the clarification of which ModelOps practices act as the primary catalysts for trust. The dominance of Monitoring, Quality Assurance, and Human oversight suggests that stakeholder trust is specifically reactive to "active" safeguards. While technical hygiene factors like versioning and auditability are foundational to the engineering pipeline, they appear to be perceived by stakeholders as "back-office" prerequisites rather than active trust-builders. This implies that building trust requires "Transparent ModelOps" where the evidence of monitoring and human control is made visible to the broader organisation.

5.2 The Behavioural Primacy of Trust

The most striking result is the nearly linear relationship between Trust and AI Usage ($\beta = .912$). This suggests that in the high-stakes, probabilistic environment of AI, trust is not merely a "nice-to-have" but the absolute gatekeeper of behavioural adoption. The disparity between trust's effect on Usage (.912) versus Perceived Value (.440) highlights a critical distinction: stakeholders may acknowledge the abstract value of an AI tool with only moderate confidence, but they will not integrate it into their daily workflows without verified trust. Trust effectively acts as the filter that transforms latent value perceptions into active usage.

The confirmation of Full Mediation (H4a and H4b) underscores the limitation of purely technical

investments. Historically, ModelOps has been viewed as an engineering optimisation problem. However, our data indicates that technical excellence has zero direct effect on usage and its entire influence is mediated by trust. If an organisation invests in automated deployment and drift management but fails to translate these safeguards into a transparent narrative for end-users, the "Technical-Behavioural Chasm" will remain uncrossed.

6. Future Implications & Contributions

The results suggest that organisations should rethink how ModelOps is implemented and communicated. Rather than treating ModelOps as a back-end engineering function, it should be positioned as a visible governance framework that actively contributes to stakeholder confidence.

Specifically, practitioners should consider the following:

- **Implement "Transparent ModelOps" Interfaces:** Since monitoring and quality were found to be the strongest drivers of trust, organisations should move dashboards out of the engineering silo. Providing business stakeholders with accessible, non-technical views of system health and performance can convert technical reliability into visible trustworthiness.
- **Formalise and Communicate Human Oversight:** Given the importance of Human-in-the-loop (HIL) approval in the measurement model, trust is clearly tied to the perception of human agency. Organisations should explicitly define and share their protocols and expert review processes to reassure users that the AI system does not operate without accountability.
- **Focus on Adoption, Not Just ROI:** The study reveals that "Perceived Value" does not automatically lead to "Usage". Management should pivot from marketing the ROI of AI tools to focusing on risk-reduction through ModelOps. If stakeholders trust the system's operational stability, usage will follow as a natural behavioural outcome.
- **Prioritise Real-Time Safeguards:** While long-term auditability is necessary for compliance, trust is built on real-time reliability. Prioritising investments in drift detection and continuous monitoring provides the immediate feedback stakeholders need to rely on AI outputs for daily decision-making.

By making these elements visible, organisations can convert technical reliability into perceived trustworthiness, thereby enabling adoption.

6.1 Theoretical Contributions

This study contributes to the literature in several ways. First, it extends the understanding of AI adoption by integrating ModelOps as an antecedent to trust, a relationship that has received limited empirical attention. Second, it demonstrates that trust functions as a full mediator, providing a more nuanced explanation of how technical and organisational factors interact. Third, it highlights the differential impact of trust on behavioural versus perceptual outcomes, suggesting that future models of AI adoption should explicitly distinguish between these dimensions.

6.2 Path Forward: Extending Beyond Reliable ModelOps

While this study provides a focused examination of Reliable ModelOps, it also opens several avenues for future research. First, the modest model fit suggests that additional explanatory factors remain unaccounted for, indicating that ModelOps alone does not fully capture the complexity of AI adoption. Future studies should extend this framework by incorporating other technical characteristics of AI systems, such as model explainability, robustness, fairness, and transparency. These factors may independently or jointly influence trust, and in some contexts, may even rival or exceed the importance of operational reliability. Second, the role of explainability and interpretability deserves particular attention.

While ModelOps ensures that systems are well-managed and governed, it does not necessarily make their outputs understandable. In domains where decision justification is critical, explainability may act as a parallel or complementary pathway to trust. Similarly, fairness and bias mitigation mechanisms may shape trust in socially sensitive applications, where ethical considerations are as important as technical performance. Third, future research could explore contextual and organisational moderators.

The importance of ModelOps and trust may vary across industries, levels of AI maturity or user expertise. For example, technical users may value versioning and auditability more highly than non-technical stakeholders, while high-risk domains such as healthcare or finance may place greater emphasis on human oversight and accountability. Finally, longitudinal studies could provide deeper insights into how trust evolves as stakeholders gain experience with AI systems. Trust may not be static; it may strengthen or erode depending on system performance, transparency, and organisational communication.

In conclusion, this study highlights that achieving successful AI adoption requires more than building technically robust systems. It requires aligning technical reliability with human trust, and future research must continue to expand this perspective by examining the broader ecosystem of technical, regulatory, organisational and psychological factors that shape AI acceptance and use.

7. Conclusion

This study set out to examine how technical reliability in AI operations translates into organisational outcomes, and the findings point to a clear and consistent narrative: reliable ModelOps practices matter, but their impact is realised only through trust. The results demonstrate that Reliable ModelOps is a well-defined, multi-dimensional construct encompassing monitoring, quality assurance, human oversight, deployment, drift management, auditability, versioning and tooling. These practices collectively form the technical backbone of AI systems. However, their influence on adoption is not direct. Instead, they shape stakeholder perceptions, with Trust emerging as the central mechanism through which operational reliability affects both perceived value and actual usage.

Among the structural relationships, the effect of Trust on AI Usage stands out as particularly strong, underscoring the behavioural importance of confidence in AI systems. While stakeholders may recognise the potential value of AI with moderate trust, actual usage appears to require a substantially higher threshold of confidence. This distinction between perceived value and behavioural adoption is critical. It suggests that organisations cannot rely solely on demonstrating efficiency gains or performance improvements; rather, they must actively cultivate trust through visible, reliable, and accountable operational practices.

The finding of full mediation further reinforces this conclusion. Reliable ModelOps does not directly translate into adoption outcomes; it must first be internalised by stakeholders as trust. This has important implications for both research and practice. From a practical standpoint, it highlights the need to move beyond purely technical implementations of ModelOps toward more transparent and communicative approaches that make reliability visible and interpretable to end-users. From a theoretical perspective, it positions trust as a necessary bridge between engineering rigour and organisational impact.

References

- [1] Amershi, S., Begel, A., Bird, C., DeLine, R., Gall, H., Kamar, E., Nagappan, N., Nushi, B., & Zimmermann, T. (2019). Software Engineering for Machine Learning: A Case Study. 2019 IEEE/ACM 41st International Conference on Software Engineering: Software Engineering in Practice (ICSE-SEIP), 291–300. <https://doi.org/10.1109/ICSE-SEIP.2019.00042>

- [2] Bharati, S., Mondal, M. R. H., & Podder, P. (2024). A Review on Explainable Artificial Intelligence for Healthcare: Why, How, and When? *IEEE Transactions on Artificial Intelligence*, 5(4), 1429–1442. <https://doi.org/10.1109/TAI.2023.3266418>
- [3] Billeter, Y., Denzel, P., Chavarriaga, R., Forster, O., Schilling, F.-P., Brunner, S., Frischknecht-Gruber, C., Reif, M., & Weng, J. (2024). MLOps as Enabler of Trustworthy AI. 2024 11th IEEE Swiss Conference on Data Science (SDS), 37–40. <https://doi.org/10.1109/SDS60720.2024.00013>
- [4] Breck, E., Cai, S., Nielsen, E., Salib, M., & Sculley, D. (2017). The ML test score: A rubric for ML production readiness and technical debt reduction. 2017 IEEE International Conference on Big Data (Big Data), 1123–1132. <https://doi.org/10.1109/BigData.2017.8258038>
- [5] Davis, F. D. (1989). Perceived Usefulness, Perceived Ease of Use, and User Acceptance of Information Technology. *MIS Quarterly*, 13(3), 319. <https://doi.org/10.2307/249008>
- [6] Gama, J., Žliobaitė, I., Bifet, A., Pechenizkiy, M., & Bouchachia, A. (2014). A survey on concept drift adaptation. *ACM Comput. Surv.*, 46(4), 44:1–44:37. <https://doi.org/10.1145/2523813>
- [7] 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. <https://doi.org/10.1108/EBR-11-2018-0203>
- [8] Haleem, A., Javaid, M., Asim Qadri, M., Pratap Singh, R., & Suman, R. (2022). Artificial intelligence (AI) applications for marketing: A literature-based study. *International Journal of Intelligent Networks*, 3, 119–132. <https://doi.org/10.1016/j.ijin.2022.08.005>
- [9] Kreuzberger, D., Köhl, N., & Hirschl, S. (2023). Machine Learning Operations (MLOps): Overview, Definition, and Architecture. *IEEE Access*, 11, 31866–31879. <https://doi.org/10.1109/ACCESS.2023.3262138>
- [10] Lingao, L. (2024). A Feasibility Study on the Application of Artificial Intelligence on the Human Resource Practices among Manufacturing Companies in China. *Journal of Digitainability, Realism & Mastery (DREAM)*, 3(02), Article 02. <https://doi.org/10.56982/dream.v3i02.211>
- [11] Lu, J., Liu, A., Dong, F., Gu, F., Gama, J., & Zhang, G. (2019). Learning under Concept Drift: A Review. *IEEE Transactions on Knowledge and Data Engineering*, 31(12), 2346–2363. <https://doi.org/10.1109/TKDE.2018.2876857>
- [12] Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An Integrative Model of Organizational Trust. *The Academy of Management Review*, 20(3), 709–734. <https://doi.org/10.2307/258792>
- [13] Mehrotra, S., Jorge, C. C., Jonker, C. M., & Tielman, M. L. (2023). Building Appropriate Trust in AI: The Significance of Integrity-Centered Explanations. *Frontiers in Artificial Intelligence and Applications*. <https://doi.org/10.3233/faia230121>
- [14] Morley, J., Floridi, L., Kinsey, L., & Elhalal, A. (2020). From What to How: An Initial Review of Publicly Available AI Ethics Tools, Methods and Research to Translate Principles into Practices. *Science and Engineering Ethics*, 26(4), 2141–2168. <https://doi.org/10.1007/s11948-019-00165-5>
- [15] Routray, B. B. (2024). The Spectre of Generative AI Over Advertising, Marketing, and Branding. <https://doi.org/10.22541/au.170534566.63147021/v1>
- [16] Sculley, D., Holt, G., Golovin, D., Davydov, E., Phillips, T., Ebner, D., Chaudhary, V., Young, M., Crespo, J.-F., & Dennison, D. (2015). Hidden technical debt in Machine learning systems. *Proceedings*

of the 29th International Conference on Neural Information Processing Systems - Volume 2, NIPS'15, 2, 2503–2511.

- [17] Shin, D. (2021). The effects of explainability and causability on perception, trust, and acceptance: Implications for explainable AI. *International Journal of Human-Computer Studies*, 146, 102551. <https://doi.org/10.1016/j.ijhcs.2020.102551>
- [18] Shneiderman, B. (2020). Human-Centered Artificial Intelligence: Reliable, Safe & Trustworthy. *International Journal of Human-Computer Interaction*, 36(6), 495–504. <https://doi.org/10.1080/10447318.2020.1741118>
- [19] Stanley, J. C., & Dorton, S. L. (2023). Exploring Trust With the AI Incident Database. *Proceedings of the Human Factors and Ergonomics Society Annual Meeting*, 67(1), 489–494. <https://doi.org/10.1177/21695067231198084>
- [20] Suresh, H., Lao, N., & Liccardi, I. (2020). Misplaced Trust: Measuring the Interference of Machine Learning in Human Decision-Making. *Proceedings of the 12th ACM Conference on Web Science, WebSci '20*, 315–324. <https://doi.org/10.1145/3394231.3397922>
- [21] Varshney, K. R., & Alemzadeh, H. (2017). On the Safety of Machine Learning: Cyber-Physical Systems, Decision Sciences, and Data Products. *Big Data*, 5(3), 246–255. <https://doi.org/10.1089/big.2016.0051>
- [22] Vibhuti Choubisa & Divyansh Choubisa. (2024). Towards trustworthy AI: An analysis of the relationship between explainability and trust in AI systems. *International Journal of Science and Research Archive*, 11(1), 2219–2226. <https://doi.org/10.30574/ijstra.2024.11.1.0300>
- [23] Yang, R., & Wibowo, S. (2022). User trust in artificial intelligence: A comprehensive conceptual framework. *Electronic Markets*, 32(4). <https://doi.org/10.1007/s12525-022-00592-6>
- [24] Zaharia, M., Chen, A., Davidson, A., Ghodsi, A., Hong, S., Konwinski, A., Murching, S., Nykodym, T., Ogilvie, P., Parkhe, M., Xie, F., & Zumar, C. (2018). Accelerating the Machine Learning Lifecycle with MLflow. *IEEE Data Eng. Bull.* <https://www.semanticscholar.org/paper/Accelerating-the-Machine-Learning-Lifecycle-with-Zaharia-Chen/b2e0b79e6f180af2e0e559f2b1faba66b2bd578a>