

Generative AI Risk Governance in Organizational Information Systems: A Conceptual Model Integrating Security, Ethics, and Digital Trust

Nurdiyanto Yusuf

Information Technology, Gunadarma University, Indonesia

Article History

Received : June 22, 2026

Revised : July 15, 2026

Accepted : July 20, 2026

Published : July 22, 2026

Corresponding author*:

Nurdiyanto.yusuf@staff.gunadarma.ac.id

Cite This Article:

Nurdiyanto Yusuf. (2026).
Generative AI Risk Governance in
Organizational Information
Systems: A Conceptual Model
Integrating Security, Ethics, and
Digital Trust. *International
Journal Science and
Technology*, 5(2), 184–202.

DOI:

<https://doi.org/10.56127/ijst.v5i2.2876>

Abstract: The rapid adoption of Generative Artificial Intelligence in organizational information systems has created new opportunities for improving productivity, decision-making, service innovation, and knowledge management. However, its implementation also introduces critical risks related to data privacy, information security, inaccurate outputs, algorithmic bias, ethical misuse, and declining user trust. **Objective:** This study aims to develop a conceptual model of Generative AI risk governance by integrating AI governance readiness, information security control, ethical AI awareness, user digital trust, and AI adoption effectiveness. The model is proposed to explain how organizations can adopt Generative AI in a secure, ethical, responsible, and trusted manner. **Methodology:** This study employed a conceptual research design using an integrative literature review approach. Data were collected from secondary academic sources, including peer-reviewed journal articles, reputable conference proceedings, and official technical reports relevant to Generative AI, information systems, cybersecurity, AI ethics, responsible AI governance, and digital trust. The data were analyzed through thematic synthesis to identify conceptual domains, relationships among constructs, and research propositions. **Findings:** The findings indicate that AI governance readiness serves as a foundational construct that strengthens information security control and ethical AI awareness. These two mechanisms contribute to user digital trust, which subsequently supports the effectiveness of Generative AI adoption in organizational information systems. **Implications:** This study implies that organizations should not adopt Generative AI solely based on technological benefits. Organizations need to establish governance policies, security controls, ethical guidelines, user education, and trust-building strategies to ensure that Generative AI implementation is safe, accountable, and aligned with organizational objectives. **Originality:** The originality of this study lies in its integrated conceptual framework, which connects technology adoption, information security, AI ethics, responsible AI governance, and digital trust into a single model for responsible Generative AI implementation in organizational information systems.

Keywords: Generative AI; AI governance; information security; ethical AI; digital trust; information systems.

INTRODUCTION

The rapid development of Generative Artificial Intelligence has transformed the way organizations design, operate, and evaluate information systems. Generative AI is increasingly used to support document automation, customer service, decision support, knowledge management, data analysis, software development, academic services, and

organizational communication. In the field of information systems, this development shows that AI is no longer positioned merely as an experimental technology, but as a strategic digital infrastructure that supports organizational productivity and service innovation (Dwivedi et al., 2021). However, the adoption of Generative AI also introduces significant risks. Large language models may generate convincing but inaccurate information, which creates challenges for information integrity and decision quality (Bender et al., 2021). In addition, Generative AI may increase risks related to data privacy, information security, bias, and human-AI interaction when it is implemented without proper governance mechanisms (Autio et al., 2024). Security threats are also becoming more complex, particularly through attacks such as indirect prompt injection that can compromise real-world applications integrated with large language models (Greshake et al., 2023). Therefore, the implementation of Generative AI in organizational information systems requires a governance-oriented approach that integrates risk management, information security, ethical awareness, and digital trust.

Previous studies related to this topic can be grouped into three major categories. First, studies on technology adoption have widely used the Technology Acceptance Model to explain user acceptance through perceived usefulness and perceived ease of use (Davis, 1989). This perspective was later expanded through the Unified Theory of Acceptance and Use of Technology, which emphasizes performance expectancy, effort expectancy, social influence, and facilitating conditions as determinants of technology use (Venkatesh et al., 2003). Further development of UTAUT also highlights the role of habit, hedonic motivation, and price value in explaining user behavior toward information technology (Venkatesh et al., 2012). Although these models provide strong foundations for understanding technology acceptance, they are not sufficient to explain the specific risks of Generative AI, such as hallucination, bias, data leakage, and accountability of AI-generated outputs.

Second, studies on information security have examined privacy protection, security awareness, policy compliance, misuse prevention, and security-conscious behavior in organizations. Bélanger and Crossler emphasize that privacy remains a central issue in information systems research (Bélanger & Crossler, 2011), particularly in relation to the collection, processing, and use of personal data. D'Arcy show that user awareness of security countermeasures can reduce information systems misuse (D'Arcy et al., 2009), while Siponen and Vance explain that employees may violate security policies through

rationalization and neutralization mechanisms (Siponen & Vance, 2010). In addition, (Ifinedo, 2012) and Safa (Safa et al., 2016) highlight the importance of behavioral and organizational factors in shaping compliance with information security policies. Nevertheless, these studies were largely developed in the context of conventional information systems and have not fully addressed emerging threats in Generative AI-based systems.

Third, studies on AI ethics, responsible AI governance, and digital trust have emphasized fairness, transparency, accountability, explainability, and trust as key elements of AI implementation. Floridi and Cowls propose ethical principles for AI in society (Floridi & Cowls, 2019), while Jobin show that transparency, justice, non-maleficence, responsibility, and privacy are frequently discussed in global AI ethics guidelines (Jobin et al., 2019). Shin further explains that explainability and causability influence user perception, trust, and acceptance of AI systems (Shin, 2021). Trust has also been identified as a critical factor in human interaction with automated and digital systems (Lee & See, 2004; McKnight et al., 2002). More recently, Yang and Wibowo developed a conceptual framework of user trust in AI (Yang & Wibowo, 2022), while Papagiannidis emphasized the need for responsible AI governance as a research agenda in information systems (Papagiannidis et al., 2025). However, existing studies still tend to discuss technology adoption, information security, AI ethics, and digital trust separately, leaving a conceptual gap in explaining how these dimensions interact in the governance of Generative AI within organizational information systems.

Based on this research gap, this study aims to develop a conceptual model of Generative AI risk governance in organizational information systems. Specifically, this study seeks to explain how AI governance readiness can strengthen information security control and ethical AI awareness, how information security control and ethical AI awareness can enhance user digital trust, and how user digital trust can contribute to the effectiveness of Generative AI adoption. This objective is grounded in the view that information system success cannot be evaluated only through system use and user satisfaction, but also through broader organizational impacts and quality dimensions (DeLone & McLean, 2003). In the context of Generative AI, adoption effectiveness should therefore be understood not only as the extent to which users accept and use the system, but also as the extent to which the system is governed securely, ethically, and responsibly. Unlike previous studies that tend to examine adoption, security, ethics, and trust as separate

research streams, this study integrates these dimensions into a single conceptual framework. The proposed model is expected to contribute to information systems literature by extending the discussion of AI adoption from a user acceptance perspective toward a governance-based perspective.

The central argument of this study is that the effectiveness of Generative AI adoption in organizational information systems is not determined solely by technological sophistication, but also by the organization's ability to govern AI-related risks. AI governance readiness is expected to influence information security control because organizations with clear policies, accountability structures, and risk management procedures are more capable of protecting data, regulating access, and mitigating security threats. This argument is consistent with the view that AI risk management requires systematic governance functions to map, measure, manage, and monitor risks throughout the AI lifecycle (Tabassi, 2023). AI governance readiness is also expected to influence ethical AI awareness because governance mechanisms can promote fairness, transparency, accountability, and responsible AI use (Floridi & Cowls, 2019; Jobin et al., 2019).

Furthermore, information security control is expected to enhance user digital trust by reducing perceived risk and increasing confidence in the reliability of AI-based systems. Ethical AI awareness is also expected to strengthen user digital trust because users are more likely to trust systems that are perceived as transparent, fair, explainable, and accountable (Shin, 2021; Yang & Wibowo, 2022). Finally, user digital trust is expected to improve AI adoption effectiveness because users are more likely to accept, rely on, and continuously use Generative AI systems when they perceive them as secure, ethical, and reliable. Accordingly, this study proposes that AI governance readiness positively influences information security control and ethical AI awareness; information security control and ethical AI awareness positively influence user digital trust; and user digital trust positively influences AI adoption effectiveness.

RESEARCH METHOD

The unit of analysis in this study is the conceptual relationship among governance, security, ethics, trust, and adoption effectiveness in the implementation of Generative Artificial Intelligence within organizational information systems. This study does not examine individual users, organizations, or technical systems as empirical objects, but focuses on theoretical constructs derived from the information systems literature. The main

constructs analyzed in this study are AI governance readiness, information security control, ethical AI awareness, user digital trust, and AI adoption effectiveness. These constructs were selected because they represent critical dimensions in understanding how organizations can manage the risks and opportunities of Generative AI. In this context, Generative AI is treated as a digital innovation embedded in organizational information systems, while governance, security, ethics, and trust are treated as conceptual mechanisms that may influence the effectiveness of AI adoption.

As illustrated in Figure 1, the research process began with the identification of the conceptual problem related to Generative AI risk governance in organizational information systems. The next stage involved selecting relevant literature from the domains of technology adoption, information security, AI ethics, responsible AI governance, digital trust, and information systems success. The selected literature was then classified into five conceptual domains, followed by thematic synthesis to identify theoretical relationships among the constructs. Finally, the results of the synthesis were used to develop the proposed conceptual model and research propositions.

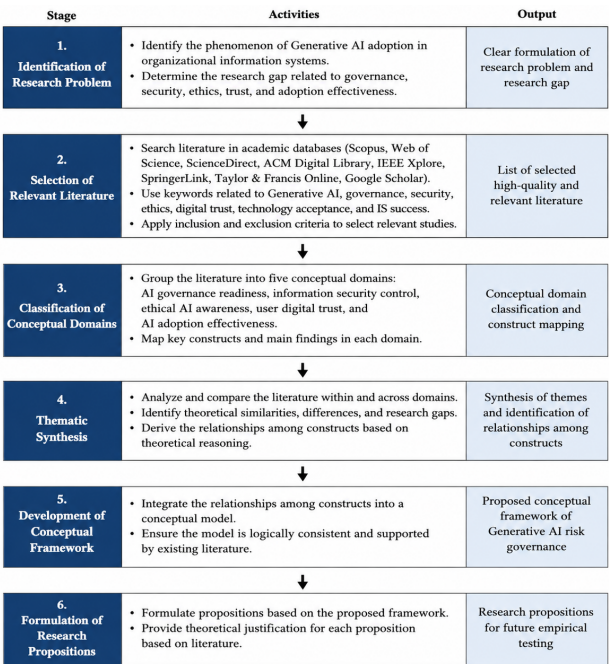


Figure 1 Research Method Flow

This study uses a conceptual research design with an integrative literature review approach. A conceptual design was chosen because the main purpose of this study is to develop a theoretical framework rather than to test empirical hypotheses using primary data. Conceptual research is appropriate when a study seeks to clarify theoretical

constructs, synthesize fragmented knowledge, and propose new relationships among concepts (Whetten, 1989). The integrative literature review approach was also used because the topic of Generative AI risk governance intersects with several research domains, including information systems, cybersecurity, technology adoption, AI ethics, responsible AI governance, and digital trust. An integrative review enables researchers to combine different theoretical perspectives and develop a more comprehensive understanding of an emerging phenomenon (Snyder, 2019). Therefore, this method is suitable for developing a conceptual model that explains the role of governance, security, ethics, and trust in the adoption of Generative AI.

The sources of information in this study consist of secondary academic literature and technical reports relevant to Generative AI and organizational information systems. The literature includes peer-reviewed journal articles, reputable conference proceedings, and official technical reports that discuss Generative AI, Large Language Models, AI governance, information security, ethical AI, digital trust, technology acceptance, and information systems success. Seminal studies were included to provide theoretical foundations, particularly those related to technology acceptance, information systems success, information security compliance, and trust in digital systems. Recent studies were also prioritized to capture the latest developments in Generative AI and AI risk governance. To maintain academic credibility, the selected sources were limited to references that have clear publication information and valid DOI numbers, except when official technical reports were used as supporting references for AI risk management frameworks.

Data collection was conducted through document-based literature identification and selection. The literature search focused on academic databases and publisher platforms such as Scopus, Web of Science, ScienceDirect, ACM Digital Library, SpringerLink, IEEE Xplore, Taylor & Francis Online, and Google Scholar for cross-checking bibliographic information. The search process used keywords and keyword combinations such as “Generative AI governance,” “AI risk management,” “responsible AI governance,” “information security control,” “AI ethics,” “digital trust,” “technology acceptance,” “information systems success,” and “Large Language Model security.” The inclusion criteria consisted of studies that discuss at least one of the main constructs in the proposed model and provide theoretical, conceptual, or empirical insights relevant to the governance of Generative AI in organizational information systems. Meanwhile, non-academic

articles, duplicated sources, opinion-based writings without scholarly support, and references without traceable DOI information were excluded from the main reference list.

The data were analyzed using thematic synthesis. The first stage involved classifying the selected literature into five conceptual domains: AI governance readiness, information security control, ethical AI awareness, user digital trust, and AI adoption effectiveness. The second stage involved identifying the key arguments, theoretical assumptions, and findings within each domain. For example, studies on technology acceptance were used to understand adoption behavior, while studies on information security were used to explain the importance of control mechanisms and compliance in protecting organizational systems. Studies on AI ethics and responsible AI governance were used to identify principles such as fairness, transparency, accountability, explainability, and human oversight, whereas studies on digital trust were used to explain how users develop confidence in AI-based systems. The final stage involved integrating these themes into a conceptual framework and formulating research propositions. Through this process, the study developed a model that explains how AI governance readiness may strengthen information security control and ethical AI awareness, how these two dimensions may enhance user digital trust, and how user digital trust may improve the effectiveness of Generative AI adoption in organizational information systems.

RESULT

The results of this study are presented as conceptual findings derived from thematic synthesis of the selected literature. Since this study is conceptual in nature, the evidence presented in this section does not come from survey responses or interview transcripts, but from theoretical patterns, conceptual relationships, and research gaps identified in previous studies. The results are organized into three subchapters in accordance with the research objective: identifying the main conceptual domains of Generative AI risk governance, explaining the relationships among the constructs, and developing a proposed conceptual model with research propositions.

Conceptual Domains of Generative AI Risk Governance

The first result of this study is the identification of five major conceptual domains that are relevant to Generative AI risk governance in organizational information systems. These domains are AI governance readiness, information security control, ethical AI awareness,

user digital trust, and AI adoption effectiveness. The selection of these domains was based on the synthesis of literature on technology adoption, information systems success, cybersecurity, AI ethics, responsible AI governance, and digital trust. Studies on technology acceptance explain that user adoption is influenced by perceived usefulness, perceived ease of use, performance expectancy, and facilitating conditions (Davis, 1989; Venkatesh et al., 2012). However, Generative AI introduces risks that go beyond traditional technology acceptance, particularly because AI-generated outputs may be inaccurate, biased, or difficult to verify (Bender et al., 2021). Therefore, the adoption of Generative AI requires a broader governance-oriented perspective.

Table 1. Conceptual Domains of Generative AI Risk Governance

Conceptual Domain	Main Focus	Synthesized Meaning
AI governance readiness	Policies, accountability, risk management, monitoring, and organizational readiness	The organizational capacity to regulate, supervise, and manage the implementation of Generative AI responsibly.
Information security control	Data protection, access control, privacy, security awareness, and misuse prevention	The ability of organizations to protect data, systems, and users from security risks related to AI-based information systems.
Ethical awareness	AI Fairness, transparency, accountability, explainability, and responsible AI use	The awareness of ethical principles required to ensure that AI systems are used fairly, transparently, and responsibly.
User digital trust	Confidence, perceived reliability, reduced perceived risk, and willingness to rely on AI systems	The degree to which users believe that Generative AI systems are secure, reliable, transparent, and beneficial.
AI adoption effectiveness	Usefulness, system success, productivity, decision quality, and continuous use	The extent to which Generative AI adoption improves organizational processes, service quality, decision-making, and user acceptance.

Table 1 shows that Generative AI risk governance cannot be understood only from the perspective of technology adoption. Instead, it requires the integration of organizational governance, security management, ethical awareness, trust formation, and adoption outcomes. In simpler terms, organizations need not only to encourage users to adopt Generative AI, but also to ensure that the system is governed, secured, ethically controlled, and trusted.

Several patterns emerge from this conceptual mapping. First, AI governance readiness functions as the starting point because governance provides the formal basis for policies, accountability, and risk control. Second, information security control and ethical AI awareness represent two important mechanisms that translate governance into operational

practice. Third, user digital trust appears as a bridge between organizational control and adoption effectiveness. Fourth, AI adoption effectiveness is not only related to usage behavior, but also to whether Generative AI improves productivity, service quality, and decision-making in a responsible manner. These patterns indicate that Generative AI adoption should be viewed as a governance-based process rather than merely as a technical implementation.

The interpretation of this finding is that Generative AI risk governance requires a multidimensional approach. Traditional technology adoption models are useful for explaining user acceptance, but they are insufficient to address the specific risks of Generative AI. The inclusion of governance, security, ethics, and trust provides a more comprehensive explanation of how organizations can adopt Generative AI responsibly. This finding supports the argument that Generative AI adoption in organizational information systems should be evaluated not only through ease of use or usefulness, but also through risk readiness, ethical responsibility, security control, and user trust.

Relationships among Governance, Security, Ethics, and Digital Trust

The second result of this study is the identification of conceptual relationships among the five domains. The literature indicates that AI governance readiness is closely related to information security control and ethical AI awareness. AI governance provides the organizational foundation for defining policies, assigning responsibilities, identifying risks, and monitoring the use of AI systems. The AI Risk Management Framework emphasizes that AI risks should be governed through structured processes of mapping, measuring, managing, and monitoring risks across the AI lifecycle (Tabassi, 2023). In the context of Generative AI, governance is also needed because the technology may create risks related to privacy, information security, bias, harmful content, and human-AI interaction (Autio et al., 2024). Therefore, governance readiness can be interpreted as an antecedent of both security control and ethical awareness.

Table 2. Synthesized Relationships among Constructs

Relationship	Conceptual Evidence	Restatement	Temporary Conclusion
AI governance readiness → Information security control	AI risk management requires policies, monitoring, access control, and risk mitigation (Tabassi, 2023). Security awareness and compliance	Organizations with stronger AI governance are more likely to	Governance readiness strengthens the organization’s ability to protect AI-based information systems.

	reduce misuse of information systems (D'Arcy et al., 2009; Siponen & Vance, 2010)	implement better security controls.	
AI governance readiness → Ethical AI awareness	Responsible AI governance emphasizes accountability, transparency, fairness, and human oversight (Floridi & Cowls, 2019; Papagiannidis et al., 2025)	Governance mechanisms can encourage ethical awareness in AI development and use.	AI governance is a foundation for responsible and ethical AI implementation.
Information security control → User digital trust	Privacy and security protection influence user confidence in digital systems (Bélanger & Crossler, 2011; McKnight et al., 2002)	Users are more likely to trust systems that protect their data and reduce security risks.	Security control contributes to stronger user trust.
Ethical AI awareness → User digital trust	Explainability and fairness influence trust and acceptance of AI systems (Shin, 2021; Yang & Wibowo, 2022)	Users trust AI systems more when they perceive them as transparent, fair, and accountable.	Ethical AI awareness strengthens trust in Generative AI systems.
User digital trust → AI adoption effectiveness	Trust influences reliance on automated systems and user acceptance of digital technology (Lee & See, 2004; Venkatesh et al., 2012)	Users are more willing to adopt and continue using AI systems when they trust them.	Digital trust supports effective and sustainable AI adoption.

Table 2 restates the conceptual relationships in a simpler form. AI governance readiness is positioned as the initial driver because it enables organizations to develop security mechanisms and ethical principles. Information security control and ethical AI awareness then contribute to the formation of user digital trust. Finally, user digital trust supports the effectiveness of Generative AI adoption by increasing user willingness to accept, rely on, and continuously use AI-based systems.

Three important tendencies can be identified from these relationships. First, governance has a dual role: it supports technical control through information security and normative control through ethical AI awareness. Second, security and ethics operate as trust-building mechanisms because users are more likely to trust AI systems when they perceive them as secure, fair, transparent, and accountable. Third, trust is not merely a psychological factor, but a strategic condition that influences the effectiveness of technology adoption. Fourth, adoption effectiveness depends on both organizational readiness and user confidence. This means that successful Generative AI implementation requires alignment between organizational governance and user perception.

The interpretation of this finding is that digital trust plays a mediating conceptual role in the adoption of Generative AI. Organizations may provide advanced AI systems, but adoption may remain weak when users perceive the systems as insecure, biased, opaque, or unreliable. Conversely, when organizations establish strong governance, implement security controls, and promote ethical AI awareness, users are more likely to develop trust in the system. This trust then becomes an important mechanism for improving adoption effectiveness. Therefore, the relationship among governance, security, ethics, trust, and adoption effectiveness forms the core logic of the proposed conceptual model.

Proposed Conceptual Model and Research Propositions

The third result of this study is the development of a proposed conceptual model of Generative AI risk governance in organizational information systems. The model integrates five main constructs: AI governance readiness, information security control, ethical AI awareness, user digital trust, and AI adoption effectiveness. The model was developed based on the thematic synthesis of previous studies and the theoretical need to integrate fragmented perspectives on technology adoption, cybersecurity, AI ethics, responsible AI governance, and digital trust. Unlike conventional adoption models that focus mainly on user acceptance, this model places governance as the starting point of Generative AI adoption. This position is important because Generative AI systems require clear governance mechanisms to manage risks related to security, ethics, privacy, accountability, and trust.

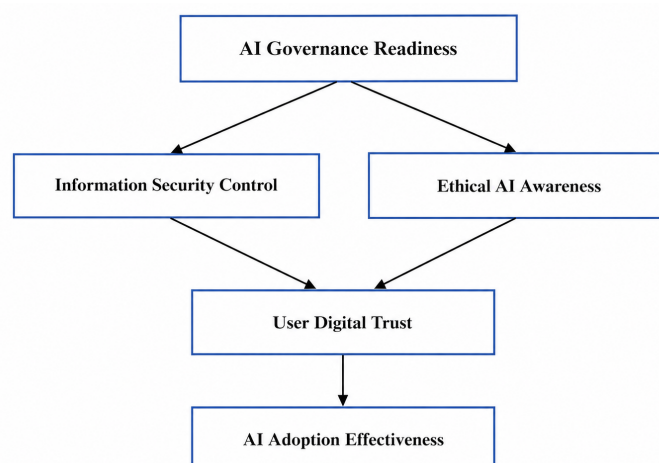


Figure 2. Proposed Conceptual Model of Generative AI Risk Governance

The model in Figure 2 shows that AI governance readiness is expected to influence information security control and ethical AI awareness. Information security control and ethical AI awareness are then expected to influence user digital trust. Finally, user digital trust is expected to influence AI adoption effectiveness. In simple terms, the model suggests that Generative AI can be adopted effectively when organizations are able to govern AI risks, secure the system, ensure ethical use, and build user trust.

Table 3. Research Propositions Derived from the Conceptual Model

Proposition	Statement	Theoretical Justification
P1	AI governance readiness positively influences information security control.	AI governance provides policies, procedures, accountability, and monitoring mechanisms needed to manage AI-related security risks.
P2	AI governance readiness positively influences ethical AI awareness.	Responsible governance encourages fairness, transparency, accountability, explainability, and human oversight in AI implementation.
P3	Information security control positively influences user digital trust.	Users are more likely to trust AI-based systems when they perceive that data, access, and system processes are protected.
P4	Ethical AI awareness positively influences user digital trust.	Users are more likely to trust AI systems that are perceived as fair, transparent, explainable, and accountable.
P5	User digital trust positively influences AI adoption effectiveness.	Trust increases users' willingness to accept, rely on, and continuously use Generative AI systems in organizational processes.

Several conceptual patterns can be drawn from Figure 2 and Table 3. First, AI governance readiness is positioned as the antecedent construct because governance determines the direction, rules, and accountability of Generative AI implementation. Second, information security control and ethical AI awareness are positioned as intermediate mechanisms because they translate governance readiness into practical and normative safeguards. Third, user digital trust is positioned as a key mechanism that connects organizational safeguards with adoption effectiveness. Fourth, AI adoption effectiveness is positioned as the final outcome because the success of Generative AI should be reflected in secure, ethical, trusted, and productive system use.

The interpretation of this model is that Generative AI adoption should not be treated as a linear process of technology acceptance, but as a governance-based process involving multiple layers of organizational readiness and user perception. The proposed model contributes to information systems literature by extending adoption theory with governance, security, ethics, and trust dimensions. It also provides practical implications for organizations. Before implementing Generative AI, organizations should develop governance policies, strengthen information security controls, promote ethical AI awareness, and build user trust. Without these elements, Generative AI adoption may produce operational benefits but also increase the risk of data misuse, biased outputs, inaccurate decisions, and user resistance. Therefore, the proposed model offers a conceptual foundation for future empirical studies on responsible and effective Generative AI adoption in organizational information systems.

DISCUSSION

The results of this study show that Generative AI risk governance in organizational information systems can be understood through five interconnected conceptual domains: AI governance readiness, information security control, ethical AI awareness, user digital trust, and AI adoption effectiveness. The first finding identifies these five domains as the main elements required to explain responsible Generative AI adoption. The second finding shows that AI governance readiness functions as an antecedent that supports both information security control and ethical AI awareness. The third finding proposes that information security control and ethical AI awareness strengthen user digital trust, which subsequently contributes to AI adoption effectiveness. These findings indicate that Generative AI adoption should not be understood merely as a matter of technological acceptance, but as a broader organizational process involving governance, risk control, ethical responsibility, and trust formation.

The relationship among these constructs can be explained by the nature of Generative AI itself. Unlike conventional information systems, Generative AI can produce new content, generate recommendations, support decision-making, and interact with users in a conversational manner. This capability creates benefits, but it also introduces risks related to inaccurate outputs, hallucination, privacy exposure, bias, and misuse. Bender explain that large language models may produce fluent but unreliable outputs, which makes information integrity a critical issue ([Bender et al., 2021](#)). For this reason, AI governance

readiness becomes important because it provides organizational rules, responsibilities, monitoring mechanisms, and risk management procedures. Tabassi emphasizes that AI risk management requires systematic governance through mapping, measuring, managing, and monitoring risks. In this study, governance readiness is interpreted as the organizational foundation that enables security and ethical mechanisms to operate effectively (Tabassi, 2023).

The finding that AI governance readiness influences information security control is consistent with previous studies in information security. D'Arcy show that awareness of security countermeasures can reduce information systems misuse (D'Arcy et al., 2009), while Siponen and Vance argue that security policy violations may occur when employees rationalize unsafe behavior (Siponen & Vance, 2010). Ifinedo also highlights that compliance with information security policies depends on both behavioral and organizational factors (Ifinedo, 2012). In the context of Generative AI, these insights remain relevant but need to be extended. Security control is no longer limited to password protection, access rights, or data storage, but also includes the protection of prompts, training data, generated outputs, model integration, and user interactions. Therefore, this study expands information security literature by positioning security control as part of a broader AI governance structure.

The relationship between AI governance readiness and ethical AI awareness is also supported by the literature on responsible AI. Floridi and Cowls argue that AI development should be guided by principles such as beneficence, non-maleficence, autonomy, justice, and explicability (Floridi & Cowls, 2019). Jobin further show that transparency, fairness, responsibility, and privacy are recurring principles in global AI ethics guidelines (Jobin et al., 2019). These ethical principles require organizational mechanisms to become practical and operational. In other words, ethical awareness does not emerge automatically from technology use. It needs to be supported by governance policies, user education, accountability structures, and clear procedures for responsible AI use. This study therefore contributes to responsible AI literature by explaining that ethical AI awareness should be treated not only as an individual moral concern, but also as an organizational capability shaped by governance readiness.

The role of user digital trust in this model is particularly important. Trust has long been recognized as a key factor in user interaction with digital and automated systems. McKnight explain that trust influences user willingness to depend on digital systems

(McKnight et al., 2002), while Lee and See emphasize that trust is essential for appropriate reliance on automation (Lee & See, 2004). In AI-based systems, trust becomes even more complex because users may not fully understand how the system generates outputs. Shin shows that explainability and causability influence user perception, trust, and acceptance of AI systems (Shin, 2021). Similarly, Yang and Wibowo argue that user trust in AI is shaped by multiple factors, including transparency, reliability, privacy, and perceived risk (Yang & Wibowo, 2022). Based on these studies, the present conceptual model positions user digital trust as a bridge between organizational safeguards and adoption effectiveness. This means that security and ethics are not only technical or normative issues, but also trust-building mechanisms.

Compared with previous technology adoption models, this study offers a broader explanation of AI adoption effectiveness. The Technology Acceptance Model explains adoption through perceived usefulness and perceived ease of use (Davis, 1989). UTAUT expands this explanation by including performance expectancy, effort expectancy, social influence, and facilitating conditions. These models remain useful for explaining why users accept technology, but they are not sufficient to explain the risks and governance needs of Generative AI. In addition, the DeLone and McLean model emphasizes system quality, information quality, service quality, use, user satisfaction, and net benefits as dimensions of information systems success. This study extends these perspectives by arguing that AI adoption effectiveness should also include governance quality, security control, ethical responsibility, and user digital trust. The novelty of this study lies in integrating adoption theory, information security, AI ethics, responsible AI governance, and digital trust into a single conceptual model for Generative AI-based organizational information systems.

The interpretation of these findings suggests that Generative AI adoption has broader social and organizational implications. In organizational settings, the use of Generative AI may influence how employees produce information, make decisions, communicate with users, and evaluate knowledge. If Generative AI is implemented without proper governance, it may create dependency on automated outputs, weaken human verification, and increase the risk of inaccurate or biased decisions. However, when it is supported by governance readiness, security control, ethical awareness, and trust-building mechanisms, Generative AI can become a responsible digital innovation that improves productivity and service quality. Thus, the proposed model contributes to a more balanced understanding of

AI adoption by recognizing both its transformative potential and its governance-related risks.

The reflection from this study shows that the proposed model has both functional and dysfunctional implications. Functionally, the model can help organizations understand the key elements required to implement Generative AI responsibly. It encourages organizations to design governance policies, strengthen security controls, promote ethical awareness, and build user trust before expanding AI adoption. This can reduce the risk of data misuse, biased outputs, user resistance, and unreliable decision-making. However, the model also implies several possible challenges. Excessive governance may slow innovation if organizations focus too heavily on control and compliance. Strict security requirements may also reduce usability if they are not designed proportionally. In addition, ethical guidelines may become symbolic if they are not translated into practical procedures and measurable responsibilities. Therefore, organizations need to balance innovation, control, ethics, and usability when implementing Generative AI.

Based on these findings, several practical actions are recommended. First, organizations should establish a clear Generative AI governance policy that defines the scope of use, user responsibilities, data protection rules, output verification procedures, and accountability mechanisms. Second, organizations should implement information security controls specifically designed for AI-based systems, including access control, data classification, prompt management, monitoring of AI outputs, and protection of sensitive information. Third, organizations should develop ethical AI guidelines that address fairness, transparency, explainability, accountability, and human oversight. Fourth, organizations should provide AI literacy and security awareness training for users so that they can understand both the benefits and risks of Generative AI. Fifth, organizations should evaluate AI adoption effectiveness not only through usage frequency or productivity improvement, but also through user trust, perceived security, ethical compliance, and decision quality. These actions can help organizations adopt Generative AI in a way that is secure, ethical, trusted, and aligned with organizational goals.

CONCLUSION

This study concludes that the effective adoption of Generative Artificial Intelligence in organizational information systems cannot be separated from governance, security, ethics, and trust. The main finding of this study shows that Generative AI risk governance

can be understood through five interconnected constructs: AI governance readiness, information security control, ethical AI awareness, user digital trust, and AI adoption effectiveness. AI governance readiness serves as the foundational construct because it provides organizational direction, policies, accountability mechanisms, and risk management procedures. Information security control and ethical AI awareness function as important mechanisms that translate governance readiness into practical safeguards. User digital trust then becomes a central bridge between organizational safeguards and the effectiveness of Generative AI adoption. Therefore, Generative AI adoption should not be viewed merely as a process of accepting new technology, but as a governance-based process that requires security protection, ethical responsibility, and user confidence.

The scientific contribution of this study lies in the development of an integrated conceptual model for Generative AI risk governance in organizational information systems. Previous studies have generally discussed technology adoption, information security, AI ethics, responsible AI governance, and digital trust as separate research streams. This study contributes by integrating these perspectives into a single conceptual framework that explains how governance readiness may influence security control and ethical awareness, how these two dimensions may strengthen user digital trust, and how digital trust may improve AI adoption effectiveness. The proposed model extends the discussion of AI adoption beyond perceived usefulness, ease of use, and system performance by incorporating governance quality, risk control, ethical principles, and trust formation. Thus, this study provides a theoretical foundation for future empirical research and offers a practical reference for organizations seeking to implement Generative AI in a secure, ethical, responsible, and trusted manner.

Despite its contribution, this study has several limitations. First, this study is conceptual in nature and does not use primary empirical data from users, organizations, or AI system implementation cases. As a result, the proposed relationships among the constructs have not yet been statistically tested. Second, the model was developed based on literature synthesis, so its applicability may vary across organizational contexts, industries, regulatory environments, and levels of digital maturity. Third, this study focuses on organizational information systems in general and does not examine specific sectors such as education, healthcare, finance, public administration, or manufacturing, where Generative AI risks may appear in different forms. Future research is therefore recommended to test the proposed model empirically using quantitative methods such as

survey-based structural equation modeling or partial least squares structural equation modeling. Future studies may also use qualitative case studies or mixed-method approaches to examine how organizations actually develop AI governance policies, implement security controls, promote ethical AI awareness, and build user digital trust in real implementation settings. By doing so, future research can strengthen, refine, or extend the conceptual model proposed in this study.

REFERENCES

- Autio, C., Schwartz, R., Dunietz, J., Jain, S., Stanley, M., Tabassi, E., Hall, P., & Roberts, K. (2024). *Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile* (NIST AI 600-1, Issue. <https://doi.org/10.6028/NIST.AI.600-1>
- Bélanger, F., & Crossler, R. E. (2011). Privacy in the Digital Age: A Review of Information Privacy Research in Information Systems. *MIS Quarterly*, 35(4), 1017-1042. <https://doi.org/10.2307/41409971>
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency,
- D'Arcy, J., Hovav, A., & Galletta, D. (2009). User Awareness of Security Countermeasures and Its Impact on Information Systems Misuse: A Deterrence Approach. *Information Systems Research*, 20(1), 79-98. <https://doi.org/10.1287/isre.1070.0160>
- Davis, F. D. (1989). Perceived Usefulness, Perceived Ease of Use, and User Acceptance of Information Technology. *MIS Quarterly*, 13(3), 319-340. <https://doi.org/10.2307/249008>
- DeLone, W. H., & McLean, E. R. (2003). The DeLone and McLean Model of Information Systems Success: A Ten-Year Update. *Journal of Management Information Systems*, 19(4), 9-30. <https://doi.org/10.1080/07421222.2003.11045748>
- Dwivedi, Y. K., Hughes, L., Ismagilova, E., Aarts, G., Coombs, C., Crick, T., Duan, Y., Dwivedi, R., Edwards, J. S., Eirug, A., Galanos, V., Ilavarasan, P. V., Janssen, M., Jones, P., Kar, A. K., Kizgin, H., Kronemann, B., Lal, B., Lucini, B.,...Williams, M. D. (2021). Artificial Intelligence (AI): Multidisciplinary Perspectives on Emerging Challenges, Opportunities, and Agenda for Research, Practice and Policy. *International Journal of Information Management*, 57, 101994. <https://doi.org/10.1016/j.ijinfomgt.2019.08.002>
- Floridi, L., & Cowls, J. (2019). A Unified Framework of Five Principles for AI in Society. *Harvard Data Science Review*, 1(1). <https://doi.org/10.1162/99608f92.8cd550d1>
- Greshake, K., Abdelnabi, S., Mishra, S., Endres, C., Holz, T., & Fritz, M. (2023). Not What You've Signed Up For: Compromising Real-World LLM-Integrated Applications with Indirect Prompt Injection. Proceedings of the 16th ACM Workshop on Artificial Intelligence and Security,

- Ifinedo, P. (2012). Understanding Information Systems Security Policy Compliance: An Integration of the Theory of Planned Behavior and the Protection Motivation Theory. *Computers & Security*, 31(1), 83-95. <https://doi.org/10.1016/j.cose.2011.10.007>
- Jobin, A., Ienca, M., & Vayena, E. (2019). The Global Landscape of AI Ethics Guidelines. *Nature Machine Intelligence*, 1, 389-399. <https://doi.org/10.1038/s42256-019-0088-2>
- Lee, J. D., & See, K. A. (2004). Trust in Automation: Designing for Appropriate Reliance. *Human Factors*, 46(1), 50-80. https://doi.org/10.1518/hfes.46.1.50_30392
- McKnight, D. H., Choudhury, V., & Kacmar, C. (2002). Developing and Validating Trust Measures for E-Commerce: An Integrative Typology. *Information Systems Research*, 13(3), 334-359. <https://doi.org/10.1287/isre.13.3.334.81>
- Papagiannidis, E., Mikalef, P., & Conboy, K. (2025). Responsible Artificial Intelligence Governance: A Review and Research Framework. *The Journal of Strategic Information Systems*, 34(2), 101885. <https://doi.org/10.1016/j.jsis.2024.101885>
- Safa, N. S., von Solms, R., & Furnell, S. (2016). Information Security Policy Compliance Model in Organizations. *Computers & Security*, 56, 70-82. <https://doi.org/10.1016/j.cose.2015.10.006>
- 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>
- Siponen, M., & Vance, A. (2010). Neutralization: New Insights into the Problem of Employee Information Systems Security Policy Violations. *MIS Quarterly*, 34(3), 487-502. <https://doi.org/10.2307/25750688>
- Snyder, H. (2019). Literature Review as a Research Methodology: An Overview and Guidelines. *Journal of Business Research*, 104, 333-339. <https://doi.org/10.1016/j.jbusres.2019.07.039>
- Tabassi, E. (2023). *Artificial Intelligence Risk Management Framework (AI RMF 1.0)* (NIST AI 100-1, Issue. <https://doi.org/10.6028/NIST.AI.100-1>
- Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User Acceptance of Information Technology: Toward a Unified View. *MIS Quarterly*, 27(3), 425-478. <https://doi.org/10.2307/30036540>
- Venkatesh, V., Thong, J. Y. L., & Xu, X. (2012). Consumer Acceptance and Use of Information Technology: Extending the Unified Theory of Acceptance and Use of Technology. *MIS Quarterly*, 36(1), 157-178. <https://doi.org/10.2307/41410412>
- Whetten, D. A. (1989). What Constitutes a Theoretical Contribution? *Academy of Management Review*, 14(4), 490-495. <https://doi.org/10.5465/amr.1989.4308371>
- Yang, R., & Wibowo, S. (2022). User Trust in Artificial Intelligence: A Comprehensive Conceptual Framework. *Electronic Markets*, 32, 2053-2077. <https://doi.org/10.1007/s12525-022-00592-6>