

An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG)

Raheela Batool ^{1,✉}, Muhammad Azlaan Zubair ²

1. *Business Administration, Pusan National University, Busan, KOR*

2. *Computer and Emerging Sciences, National University of Computer and Emerging Sciences, Karachi, PAK*

Received: January 31, 2026 | Review began: February 22, 2026 | Review ended: September 07, 2026 | Published: September 14, 2026

© Copyright 2026

This is an open access article distributed under the terms of the Creative Commons Attribution License CC-BY 4.0., which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

Abstract

The growing integration of artificial intelligence (AI) into consumer products has raised significant concerns regarding privacy, security, and user control. Despite these concerns, limited research has explored how consumers perceive and evaluate the risks associated with AI systems. This study examines three key AI-related risks: informational intrusion, cyber exposure, and loss of control. It argues that these risks are largely driven by a central factor, algorithmic opacity, which refers to the lack of transparency and the difficulty consumers face in understanding how AI systems operate.

Data for this study were collected through online communities, including the Pakistan Students and Scholars Association South Korea, which has approximately 8,000 followers, and Seekers, a diverse network comprising students, professionals, businesspeople, laborers, artists, and other community members, with approximately 17,000 followers across both platforms. Using a structured questionnaire, a total of 300 valid responses were gathered and analyzed using SPSS and AMOS 26.

The findings reveal that all three identified risks are strongly associated with perceptions of algorithmic opacity. In turn, perceived opacity is significantly linked to feelings of unfairness and discomfort among consumers. By contrast, the three risks exert only weak direct effects on these outcomes, suggesting that consumers respond more strongly to the transparency of AI systems than to the risks themselves. Overall, the results indicate that the relationship between governance-related AI risks and negative consumer responses is largely mediated by algorithmic opacity. In other words, a lack of transparency serves as the primary mechanism through which AI-related risks influence perceptions of unfairness and discomfort.

From a practical perspective, the findings highlight the importance of improving transparency in AI systems. Organizations can mitigate negative consumer perceptions by providing clearer explanations of AI-driven decisions and offering users greater control over AI-enabled processes. Enhancing transparency and understanding can increase perceptions of fairness and comfort, ultimately strengthening consumer trust in AI systems. Therefore, improving the transparency, explainability, and user comprehension of AI technologies should be a key priority for organizations seeking to foster greater acceptance and trust in AI-driven products and services.

Categories: Computational Intelligence and Information Management, Consumer Behavior, Risk Management

Keywords: artificial intelligence, information governance, algorithmic opacity, consumer risk perception, value-based adoption model, algorithmic inequity

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

Introduction

The widespread adoption of artificial intelligence (AI) in consumer-facing industries such as the healthcare sector, the finance sector, and the electronic commerce sector has revolutionized the way data are processed and utilized. However, the adoption of AI is also related to various risks and concerns, such as privacy, security, and transparency (Olawade et al., 2025) (Bulgurcu et al., 2010). Yet, the literature has paid more attention to the risks and concerns of AI from organizational and regulatory perspectives without considering the dimension of consumers. This is an important dimension to be taken into consideration because the perception of risks by consumers is a critical dimension in the development of trust in AI. The conventional literature on information governance has largely focused on organizational and regulatory mechanisms, including data governance frameworks, Information Technology (IT) governance structures, data quality management, and compliance with regulations such as General Data Protection Regulation. The importance of information management by an organization through formal policies and infrastructures of control is also emphasized by this research. However, AI governance goes beyond these organizational structures by involving consumers as stakeholders in the AI governance equation. Thus, the adoption and assessment of AI are no longer dependent on formal structures of compliance but are instead dependent on consumer concerns such as data intrusion, cyberattacks, lack of control, and a lack of transparency and fairness in decision-making processes (Joshi et al., 2025). However, the governance of AI goes beyond institutional mechanisms and places consumers as active stakeholders whose risk perceptions and benefits have a direct impact on governance (Beerends and Aydin, 2025). The adoption of AI information governance by consumers is measured by factors such as data intrusion, cyber threats, loss of control, and the lack of transparency and equity in algorithmic decision-making processes, as opposed to legislative compliance (Potluri, 2025). The study will use the Value Based Adoption Model (VAM) by Kim et al. and Alreemy et al. to offer theoretical support for the perceptions (Kim et al., 2007) (Alreemy et al., 2016). Unlike traditional theories of technology adoption, the VAM considers customers to be rational and value-driven decision-makers who weigh the perceived benefits of a technology against the perceived costs before adoption (Schneider, 2024). The “perceived sacrifice” part of perceived value is where most of the problems associated with AI technology lie (Nexhip et al., 2025) (Weber, 2020).

In order to ground the article theoretically, this study adopts the VAM in the context of AI governance. VAM is a framework for understanding the adoption of technologies as a trade-off between perceived benefits and perceived sacrifices. In AI-enabled settings, governance-related factors such as privacy risks, data security threats, lack of transparency, and reduced user control are considered critical sacrifices in AI-enabled settings. These factors are dominant concepts in AI governance research, especially when discussing AI explainability, accountability, and fairness. By considering these governance factors in VAM, the current research provides a value-based approach to consumer evaluation of AI systems that is influenced not only by the functional benefits of AI systems but also by the governance risks of AI adoption (Lämmermann et al., 2024). In particular, the concept of algorithmic opacity is considered a critical mechanism for linking VAM with AI governance concepts such as fairness (Lin et al., 2022) (Jalali et al., 2021). CCER reflects beliefs about the dangers of data breaches and security threats from AI-enabled infrastructures. EUA represents perceptions that individual control, contestability, and accountability regarding AI systems’ processes and outputs have diminished. These governance-related sacrifices reduce consumers’ perceived value of AI applications and reveal underlying weaknesses in AI information governance (Tallon et al., 2013). A central governance implication of these risks is AOR, defined as the perception that AI systems operate as “black boxes” whose decision rationales, data-processing mechanisms, and accountability structures cannot be clearly understood. The opacity of the algorithm hinders the user’s sense-making capability, which has an effect on the transparency, credibility, and trustworthiness of governance systems (Harvey and Harvey, 2014). The opacity of the algorithm is a significant sense-making process within the VAM paradigm that translates the sacrifices of governance into general assessments of value and justice. Miscommunication and misunderstandings often cause a sense of injustice in algorithmic systems, particularly when users cannot understand the processes of decision-making, risk, and allocation of responsibility (Peng et al., 2022). Notably, injustice is not only caused by unfavorable outcomes (Hammerschmidt et al., 2025). Instead, depending on the design and quality of explanation, users may view AI systems as unfair, insufficiently protective, or overly beneficial to some users (Fernández Llorca et al., 2025). These perceptions suggest that there are underlying issues of governance, in which a lack of

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

transparency and explanation mechanisms is a factor in a lack of legitimacy and user confidence (Fischer-Preßler et al., 2023) (Joia and Correia, 2018). The rapid advancement of AI technologies, especially in areas such as generative AI, automated decision-making, and personalized digital services, has also raised concerns related to data privacy, security, transparency, and fairness. The recent research on AI governance also emphasized the issues of explainability, bias, and accountability in AI systems, especially where AI plays an important role in the decision-making process of consumers in different sectors such as healthcare, finance, or e-commerce. Yet, the literature on AI governance appears to be more focused on the organizational or regulatory side of AI governance, while the consumer side has not been emphasized (Kooper et al., 2011). The research, therefore, attempts to offer an updated and relevant discussion in the backdrop of the current AI applications and AI governance. Specifically, the use of algorithmic opacity as a mediator of the relationship between the erosion of user agency and perceived inequity has not been examined (Ogbanufe et al., 2021).

To fill this research gap, this study develops and examines a value-based AI information governance framework by integrating the VAM with governance-related risk dimensions. There has been limited research on consumer risk factors such as Informational Intrusion Risk (IIR), Consumer Cyber Exposure Risk (CCER), Erosion of User Agency (EUA), Algorithmic Opacity Risk (AOR), and Perceived Algorithmic Inequity and Discomfort (PAID). This current study seeks to fill this gap by exploring the effects of these risk perceptions on how consumers evaluate AI systems. This research aims to fill the existing gap by applying a value-based AI information governance framework by combining the VAM with consumer perception of risk factors associated with AI governance. The dissertation of the author, which focused on consumer risk perceptions of AI adoption, is further developed in this manuscript. Nevertheless, this study contributes to the existing dissertation by incorporating the factors of consumer perception of AI-Driven Information Governance (AIG) and examining the direct and indirect relationships between them, which are not covered in the dissertation. Particularly, this research: Examines the direct relationships between IIR, CCER, and EUA and AOR (RQ1-RQ3). Unravels the direct associations between IIR, CCER, and EUA and perceived algorithmic inequity and algorithmic discomfort (RQ4-RQ6). Unpacks the association between AOR and PAID (RQ7). Investigates the mediating association of AOR between EUA and perceived algorithmic inequity (RQ8). This research work makes three significant contributions to the existing body of knowledge. First, it extends the existing body of knowledge on information governance in AI by adopting a consumer value and risk perception approach rather than an organizational control-based approach to information governance. Second, it situates the notion of algorithmic opacity in its rightful place as an elementary form of governance mechanism that enables risk perception and fairness assessment. Third, it validates the transparency and accountability-driven governance structures that have the potential to enhance trust and mitigate fairness concerns in AI-enabled systems. These sections comprise the remaining portions of this article. The context of the study, the governance risks of AI-based applications, and the theoretical foundation of VAM are all covered in this section.

Literature review

Information Governance

Information and governance systems have always been stressed by scholars as having strategic importance to affect organizational outcomes in different organizational settings. Tallon et al. define information governance as an asset from a strategic standpoint and illustrate that particular practices which enhance company performance are required for the effective implementation of information governance (Tallon et al., 2013). It is important to take into account the relationships between stakeholders and organizational needs when designing the governance system. Fischer-Preßler et al. (Fischer-Preßler et al., 2023) examines the notion of information governance in the case of humanitarian NGOs, highlighting that social complexity tends to exacerbate the issues related to governance that can be alleviated through the increase of social capital and process management. Although the sources focus on the factors that determine good governance internally and behaviorally, Gandía (Gandía, 2008) provides an insight into external information governance. Collectively, these articles show that despite different approaches to governance in various spheres such as corporations, humanitarian actions, and digital disclosures, its effectiveness always relies on how organizations manage information flow, align governance practices to stakeholders' needs, and adapt to emerging changes. Collectively, these articles show the importance of information governance for both society and organizations. It was found that organizations can create both big and small innovations especially in fast-changing environment if they properly manage information and utilize

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

big data. Shifting our focus from organizations to the public sector. Mansoor ([Mansoor, 2021](#)) showed that people gaining transparent information through social media can increase their faith in government and their engagement in the decision-making process if the procedures of governance are effective especially in relation to the COVID-19 problem. Both of these articles highlight the importance of accurate information, but Gaozhao provides yet another viewpoint ([Gaozhao, 2021](#)). It was found that despite the use of fact-checking symbols, the ability of people to detect any falsehood is not significantly improved with these symbols. In general, all these studies create an overall picture that goes from innovation in organizations to trust and information processing. The authors assert that good governance and information can exist only if individuals think about the information provided to them, understand it, and engage in it. Overall, this literature helps us understand data governance and information. To give the exact definition of information governance and four main research questions based on IT governance ([Donaldson and Walker, 2004](#)) provide the main concepts. Extending this theoretical framework, Abraham et al. ([Abraham et al., 2023](#)) investigate real-life data governance issues in 13 marketplaces in eight countries, categorize typical problems, and provide solutions that can be applied. Additionally, Vial explores the way businesses implement data governance practices and reveals four major topics, such as fostering innovation, managing implementation challenges, responsible data management, and service-oriented approach ([Vial, 2023](#)). Finally, the discussion goes on to cover the topic of public digital innovation initiatives in Gegenhuber et al. ([Gegenhuber et al., 2023](#)), where the authors explore the issue of distributed governance and its challenges by using the "Wir Vs. Virus" example and develop a framework based on transparency, accountability, and power. From the theoretical foundations of information governance through practical issues in marketplaces to implementation processes and broader social applications in collaborative digital projects, this literature represents an obvious connection between all the above-mentioned topics ([Abraham et al., 2019](#)). To develop a data governance framework, this study analyzes 145 sources, reveals major components, and sets directions for future research.

The structure of governance in the field of information technologies for senior care services is described, proving that institutional pressure leads to government-supported projects for better service provision, whereas market pressure affects governance mechanisms directly. Information governance in dynamic organizational networks is investigated in regard to quality, security, and metadata, and one can see the problems that arise in practice in terms of exchanging information ([Rasouli et al., 2016](#)). The contribution of resource-based information processing to competitive advantage due to more effective decision-making is proved empirically using data collected in 633 UK firms ([Cao et al., 2019](#)). Considering the health sector, the lack of properly trained HIMS specialists has been proven to have a negative impact on the practices of data management and overall performance of governmental health institutions in Malawi ([Chima et al., 2023](#)). The risks associated with health data breaches are considered from the perspective of an integrative approach providing suggestions for efficient risk management ([Pool et al., 2024](#)). Vlogs made by travelers affect individual tourists and provide information about popular behavior and cognitive trends with regard to gender differences. Risks associated with sharing patients' stories online indicate that 46.6% of stories about patients pose a risk, and 32.1% of the tweets posted could be identified by acquaintances. Crisis data management bias is noted as one of the determinants of decision-making, where debiasing techniques and awareness are suggested to mitigate the bias ([Paulus et al., 2024](#)). Organizational digital hoarding practices have been explored, and it was found that workplace culture, anxiety, and accountability result in hoarding and noncompliance with data protection policies by employees even though they are well aware of these policies. Information management problems for authors and artists in New Zealand are presented, paying attention to the long-term information management issue. In relation to an EHR system in a Canadian health authority setting, an Information Systems Utilization framework that seeks to explain the enactment of systems has been suggested ([Burton-Jones and Volkoff, 2017](#)). Consumer attitudes toward AI reveal higher acceptance when AI is presented as a complement to humans while resistance develops when AI is applied in tasks that require warmth in their interactions ([Peng et al., 2022](#)). As far as the security of protected health information is concerned, 12.3% of the population of the U.S. keeps some information from medical professionals due to their fears. Breach problems in the healthcare industry are also discussed where it is noted that the main threats include hacking and unauthorized disclosure; therefore, there is a necessity to implement more confidentiality measures ([Almulihhi et al., 2022](#)). Compliance with security among the staff members is also considered as it is determined by the personal attitude, self-efficacy, as well as by the benefits and costs of such activities.

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

Governance, and Human Agency in Sociotechnical Systems

According to academic literature, there should be an audit that would involve both documentation and replication. The necessity for defined responsibilities is considered vital when improving accountability and addressing governance problems, apart from other issues. Along with these concerns, the life-cycle model of accountability underlines the importance of distributing accountability within different parties involved in the operation of AI technology. Rather than relying solely on developers for accountability, such frameworks emphasize the necessity of placing accountability among deployers, users, managers, and organizations throughout various stages of the AI lifecycle, starting with design through deployment, utilization, monitoring, and ultimately deactivation. It is worth noting that such an approach is based on the sociotechnical aspect of AI applications whereby many of its impacts arise due to its interaction with the organizational environment in which it operates. In this regard, the concept of AI assurance becomes increasingly recognized as a governance process rather than a compliance practice. With the help of the above-stated criteria, the studies offer an actionable solution to responsible AI. The above stream of research highlights the need for risk-based AI governance frameworks to be complemented by more all-embracing and comprehensive types of AI governance frameworks that are polycentric, adaptive, and socially-oriented. Although it is right to acknowledge that there is a need for risk-based frameworks in order to assess the risks and threats associated with AI, some of the weaknesses of this type of AI governance framework include its inability to factor in the effect of the incorporation of AI on society (Mueller, 2025). It is noted further that there are some other aspects to the issue of AI governance such as international and geopolitical ones.

The processes of AI innovation and implementation take place in an environment where international competition, international fragmentation in terms of regulatory mechanisms, and differences in technology are typical. Neglecting these factors in the course of implementing certain governance policies might aggravate inequalities worldwide, increase the temptation of regulatory arbitrage, and reduce opportunities for international cooperation (Ogbanufe et al., 2021). At the organizational level, research tends to highlight the significance of internal governance controls in bringing regulatory expectations into effect. The effective management of AI-related risks requires accountability and decision-making systems to be established within the organization (Passlack et al., 2025). It is for this reason that effective AI governance should not simply depend on regulations imposed from the outside but must be complemented by efforts made by organizations in terms of how they govern themselves. Overall, AI governance emerges as an issue that encompasses many different layers (Samuel et al., 2022). The discussed research problem constitutes an interesting criticism of the ethics frameworks being static, universalistic, and principle-based as they might not be able to capture the actual practice of AI deployment. AI ethics needs to be grounded in a participatory and situational approach in order to understand the sociotechnical dynamics through which AI operates. Not only the algorithms themselves but also the whole sociotechnical dynamic, including power relations, are matters of ethics (Tiwana et al., 2013). Expanding upon this line of criticism, the unevenness in terms of responsibilities and exposure to potential harms related to the use of AI technologies among the various actors within the AI ecosystem demonstrates that marginalized groups face higher levels of risks from AI development and deployment, whereas the power to make decisions and to take actions regarding such technology is held by powerful institutions.

Information, IT, and Information Security Governance

The literature on information governance encompasses diverse domains, including information governance, IT governance, healthcare, and humanitarian organizations, providing a broad foundation for understanding governance mechanisms across different contexts (Lämmermann et al., 2024) (Nicho, 2018). Prior studies have highlighted critical themes such as transparency, trust, innovation, and socio-technical perspectives, while also documenting the evolution of governance frameworks from traditional IT governance approaches toward more complex and adaptive governance models (Cao et al., 2019). In addition, the literature has examined governance issues from behavioral, organizational, and technological perspectives, adopting an interdisciplinary approach to understanding governance effectiveness (Mikalef et al., 2020) (Mansoor, 2021). Research has further explored governance-related challenges across various organizational environments, including public institutions, healthcare systems, nongovernmental organizations, and digital platforms (Gaozhao, 2021) (Donaldson and Walker, 2004) (Abraham et al., 2023) (Gegenhuber et al., 2023). Despite these valuable

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

contributions, existing research remains dispersed across multiple contexts, with concepts such as governance, trust, information flow, privacy, and security often examined independently rather than within a unified AI governance framework (Chima et al., 2023) (Pool et al., 2024). Consequently, while substantial knowledge exists regarding governance practices and information management, limited attention has been devoted to integrating these perspectives into a coherent explanation of governance-related risks and consumer evaluations of AI-enabled systems (Almulihi et al., 2022) (Baskaran et al., 2013). Furthermore, previous studies have emphasized the importance of governance structures, security practices, and information management processes (Bulgurcu et al., 2010) (Koooper et al., 2011) (Soomro et al., 2016), yet the relationships among these dimensions remain insufficiently articulated within the emerging AI governance literature. As governance research increasingly incorporates issues related to transparency, accountability, and information quality, greater theoretical integration is needed to explain how these factors collectively influence perceptions of AI systems (De Haes and Van Grembergen, 2009) (Brown, 2006) (Xue et al., 2008) (Nicho, 2018).

Similarly, the literature on IT governance, cybersecurity, and emerging AI governance has identified numerous factors that influence governance effectiveness, including behavioral factors such as employee attitudes and self-efficacy, organizational factors such as leadership and culture, and structural factors such as governance frameworks and organizational maturity (Islam et al., 2018) (Hong et al., 2014) (AlGhamdi et al., 2020) (Kim et al., 2013) (Lin et al., 2022). Existing studies provide valuable insights into cybersecurity practices, organizational controls, healthcare-related information security concerns, and AI governance mechanisms (Hassidim et al., 2017) (Chernyshev et al., 2019) (Choi and Johnson, 2021). However, these research streams have largely developed independently, resulting in limited theoretical integration across governance, cybersecurity, and AI-related perspectives. As a result, important questions remain regarding how governance-related risks influence consumer perceptions of transparency, fairness, and accountability in AI-enabled environments (Hassandoust et al., 2021) (Novelli et al., 2024). Addressing this gap requires a more integrated framework that synthesizes behavioral, organizational, technological, and governance perspectives to better explain consumer responses to AI applications (Johnston et al., 2019).

Consumer Risk Perception in AI Applications

Recent studies have been focusing increasingly on consumer risk perceptions in the context of AI applications, and these studies have shown that people assess AI-powered systems using multi-dimensional risk matrices that go beyond traditional security concerns (Hassandoust et al., 2021). One of the most influential theoretical paradigms underpinning this stream of research has its roots in the digital health technology literature, where the process of consumer decision-making is often described using the privacy calculus framework. The empirical evidence suggests that people make assessments of AI technology by systematically comparing the expected benefits of AI technology, such as personalization and health benefits, to the potential privacy risks, such as surveillance, unauthorized access to data, and misuse of personal data (Kordzadeh et al., 2016). Risk perception is a heterogeneous construct that is shaped by demographic characteristics, health status, emotional attachment, and context, which makes risk perception a subjective and context-dependent process. In this body of research, privacy risk perception is identified as a significant predictor of behavioral intention (Majrashi, 2025). The adoption of technology, therefore, is mainly dependent on the evaluation of users as to whether the benefits are sufficient to outweigh the potential informational risks (Odilla, 2024). Apart from the risks associated with privacy, trust is a mediating variable between risk and the adoption of AI. On the basis of previous studies, trust in AI technology has been found to be significantly affected by governance-related factors such as transparency, explainability, auditability, and responsible data practices (Hammerschmidt et al., 2025). If the AI system is perceived as a “black box,” then users are bound to face more uncertainty and lack of control, which will eventually lead to an increased perception of risk. But the addition of explainability modules, accountability, and governance can help in decreasing the uncertainty and improving the perceptions of trustworthiness in the system (Recker et al., 2019). Trust, therefore, cannot be perceived as a psychological phenomenon but also as a governance-embedded phenomenon that plays a very important role in risk perception and the adoption of AI technology. At the macro-level, there are some developing models of risk perception in AI that help in understanding the significance of different types of risks from the consumer’s perspective (Novelli et al., 2024).

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

The models highlight the significance of governance factors such as safety, fairness, accountability, and risk management. Although these models are designed for organizational governance, they have an indirect relationship with how consumers perceive and classify AI systems as “high risk.” The absence of effective safety mechanisms, autonomous decision-making errors, or the absence of effective regulatory control are some factors that contribute to increased risk perception by consumers (Tiwari and Farag, 2025). Institutional norms and regulatory categorizations have a significant bearing on public perceptions about the acceptability of AI systems. Fairness and bias in AI systems further tend to heighten risks (Beerends and Aydin, 2025). Empirical evidence indicates that biased outcomes and decision-making processes of algorithms have been shown to negatively impact vulnerable groups, thereby perpetuating concerns of imbalance and damage to reputation. It is the potential for structural bias within algorithms that perpetuates concerns of imbalance and loss of autonomy, ultimately resulting in a lack of trust. From a socio-technical perspective, the narratives that societies develop around AI are significant in influencing public perceptions of intelligence, autonomy, and agency (Krokowski and Hirsch-Kreinsen, 2025). Hyperbolic stories about the potential of AI, particularly in the realm of artificial general intelligence, could potentially fuel both overenthusiasm and overconcern about loss of control or displacement by technology. Recent research also highlight the need to integrate fairness, robustness, and transparency into the development of AI. Bias correction and data management in the early stages of AI development have been identified as critical strategies to counter negative outcomes in the long term (González-Sendino et al., 2024). The comprehensive application of explainability and accountability approaches throughout various stages of system development improves transparency, increases the legitimacy of governance, and promotes greater consumer trust, thereby reducing perceptions of uncertainty (Churamani et al., 2023). Overall, the literature conceptualizes the complexity of consumer risk perceptions of AI as multidimensional and socially constructed. Privacy, security, fairness, transparency, governance quality, and ethical integrity are dynamically intertwined in the process of trust-building, which determines the eventual acceptance or rejection of AI-enabled technologies by consumers.

Hypothesis development and theoretical framework

Value Based Adoption Model (VAM)

VAM proposed by Kim et al. claimed that new ICT users should not be recognized as simply technology users, but also as consumers (Kim et al., 2007). The VAM considers the benefits and sacrifices of consumers to analyze their perceived value for new technologies such as AI applications. Specifically, the perceived value refers to the sum of the benefits the user obtains from the AI applications and the relative sacrifices the users make (Hassandoust et al., 2021). The identification of the sacrificial value can reflect consumers’ decisions in real time (Novelli et al., 2024). Value is analyzed in terms of risks and costs consumers who use AI applications pay. In order to reflect the AI applications value, consumers’ perceptions about the risks and sacrifices are composed as variables to analyze the perception of the cost that consumers had to pay to use the AI applications in preparation for consumer benefits. Referring to the factors used in the research model, novel perceived risk-related variables are introduced based on VAM constructs (AlGhamdi et al., 2020). As for the specific variables, the perceived consumers sacrifice defined as IIR, CCER, and EUA, following hypotheses are used to analyze the effect of perceived consumer AI applications information governance on AOR and AID according to the composition of the VAM model (Nicho, 2018). The VAM is a conceptual framework that assesses the balance between the benefits and sacrifices that consumers perceive in adopting new technologies (Krokowski and Hirsch-Kreinsen, 2025). In the context of AI applications Information Governance features, VAM becomes an effective tool for understanding how perceived value affects consumer perception. Therefore, in this study, the VAM is the theoretical framework because overall, the application of VAM in the context of consumer perception of AI applications provide valuable insights into how perceived value is related to AI applications features and consumers perception (Kim et al., 2013). This understanding is important for marketers and technology developers that meets consumer expectations and encourages AI applications in the online platforms. The principles of cost-benefit analyses are exemplified in the concept of value, which is broadly defined as the trade-off between total benefits received and total sacrifices (Gwebu et al., 2020). A value-based model would be able to capture the risks and sacrifices factors as a comparison of benefits and costs. We propose and empirically test a VAM of AI applications Information Governance by integrating the most relevant findings of AI applications information governance and value literature. This combined framework represents a novel approach to understanding consumers’

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

perception of AI applications information governance. Our findings should help in the theoretical understanding of AI applications information governance in a voluntary and personal context. In practice, our findings could guide AI technology developers to reduce risks and threat perceptions of AI applications.

Informational Intrusion Risk and Algorithmic Opacity in AI Information Governance

IIR is the perceived user's risk that AI applications watch, gather, alter, or manipulate personal information beyond the acceptable boundaries of privacy, control, and transparency (Tallon et al., 2013). The use of personal information in an excessive, unjustified, or poorly regulated manner leads to a perception of violation of personal information boundaries. Therefore, IIR is conceptualized as a perceived sacrifice that affects users' attitudes toward AI technologies negatively. The conceptual framework of IIR is grounded in the VAM developed by Kim et al., which contends that technology adoption is a function of the trade-off between perceived benefits and perceived sacrifices (Kim et al., 2007). Sacrifices in VAM encompass risks, potential losses, and non-financial costs involved in the adoption of technology. In the context of AI dominance, intrusive data practices are likely to increase the cost of governance due to growing concerns regarding privacy, lack of control, and unclear procedures for data processing (Odilla, 2024). With the increasing cost of governance, the user is likely to experience a lack of transparency, accountability, and process understanding in AI systems (Fernando and Dawson, 2009). Empirical evidence suggests that excessive or unjustified information controls undermine trust in governance structures as well as in transparency and accountability mechanisms (Ryan et al., 2025). Furthermore, ambiguous data flows, automated surveillance practices, and insufficiently explained system operations create unclear decision-making contexts that render algorithmic processes difficult to evaluate and interpret (Krokowski and Hirsch-Kreinsen, 2025). These concerns are particularly pronounced among populations that are disproportionately affected by algorithmic decisions, especially in sensitive domains such as healthcare (Giannouchos et al., 2021). Consequently, when governance mechanisms are perceived as ineffective in regulating and limiting informational intrusion, users are more likely to experience heightened AOR, which ultimately reduces trust in AI systems (Papagiannidis et al., 2023). The authors recognized that this manuscript adds value to the literature by highlighting the importance of transparency as an intermediary factor for risk perception reduction and trust in AI-powered information governance. Nevertheless, the study contributes to existing research by focusing on transparency as an important mechanism to lower consumer risk perceptions. The authors have tried to protect the population from disadvantages and negative aspects of AI applications usage through this small level survey study by comprehensible and scholarly synthesis and recommendations based on available data.

H1: Informational Intrusion Risk has a positive effect on AOR in AI information governance.

Consumer Cyber Exposure Risk and Algorithmic Opacity in AI Information Governance

CCER refers to the extent to which consumers feel vulnerable to cyber-attacks such as hacking and the absence of security protection while using AI applications. Consumers are also concerned about the ability of AI technology to safeguard personal and private data in a complex data environment. In the VAM, CCER is viewed as a perceived sacrifice that may lower the perceived value of an AI application. The logic of VAM is that if perceived sacrifices, such as perceived system risk, uncertainty, and losses, outweigh perceived benefits, then consumers will develop a negative value perception of the technology (Kim et al., 2007). In the context of an AI-supported environment, the higher the level of perceived cyber exposure, the greater the psychological and cognitive costs of system use, and therefore the lower the perceived value. The absence of sufficient cybersecurity and transparency in the security process may have a negative effect on the users' perception of system reliability and quality of governance, as suggested by previous studies (Choi and Johnson, 2021). Moreover, the technical complexity of AI-assisted data management makes it difficult for users to comprehend the process of mitigating cyber threats (Krokowski and Hirsch-Kreinsen, 2025). Because of the widespread use of "black box" AI systems, consumers cannot effectively monitor the detection and mitigation of data security threats, thus contributing to the perception of reduced control and lack of accountability. The absence of transparency is especially important in high-risk AI applications, such as the healthcare industry, where the lack of transparency in cybersecurity controls further exacerbates concerns about governance and accountability. From the perspective of VAM cost-benefit analysis, high CCER is indicative of governance-related shortcomings, such as a lack of transparency,

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

inadequate oversight, and poor security capabilities (Kim et al., 2007). As the level of perceived sacrifices increases, consumers are likely to perceive AI systems as opaque and ill-governed, thus contributing to high AOR. In this way, CCER acts as an indicator of VAM-based perceived risk that shapes consumer perceptions of algorithmic opacity, thus creating uncertainty about the data management and security capabilities of AI systems.

H2: Consumer Cyber Exposure Risk has a positive effect on AOR in AI information governance.

Erosion of User Agency and Algorithmic Opacity in AI Information Governance

EUA is described as the extent to which users perceive a loss of control and decision-making authority while interacting with AI systems. User agency refers to the users' ability to understand decision-making processes, locate accountability, and adequately contest or influence system outcomes. Regarding the application of AI information governance, the loss of user agency is a challenge that is important in situations where there is a lack of clarity in accountability, where accountability is shared by different actors, and where decision-making is not transparent enough (Krokowski and Hirsch-Kreinsen, 2025). In applying the VAM, the loss of user agency is a cost that is significant and reduces the value of AI applications. The VAM model argues that consumers assess technologies based on the value they can create, as well as the costs of loss of autonomy, control, and mental effort (Kim et al., 2007). In AI-assisted systems, when users do not have the capability to understand and control the system's behavior, and when accountability for the decision-making process is shared among developers, deployers, and institutions, the psychological and governance costs of using the system are significantly increased. Moreover, from the available literature, it has been observed that when there is a lack of accountability or insufficient enforcement of accountability, it results in a lack of control and understanding, thereby perpetuating the users' feelings of algorithmic opacity (Chernyshev et al., 2019). Moreover, narratives that have a tendency to overemphasize the capabilities of AI, as well as ambiguous governance structures, also add to the lack of accountability for system behavior and possible negative outcomes, making it difficult for users to contest decisions or seek redress (Choi and Johnson, 2021). Although technical and governance strategies such as audits, explainability techniques, and compliance systems are often suggested to deal with accountability, these are more institutional-level solutions and are not accessible to end users. Consequently, they are unable to fully restore individual-level agency. Based on the cost-benefit rationale of VAM, the degradation of user agency means that there is a growing gap in the value perceptions, where the perceived costs outweigh the perceived benefits in governance-related contexts (Kim et al., 2007). Due to decreased user control, recourse, and understanding, AI systems are perceived as unaccountable and opaque. Consequently, the erosion of user agency acts as a VAM-based perceived risk predictor that directly exacerbates AOR as a governance-level consequence of reduced user empowerment.

H3: Erosion of User Agency has a positive effect on AOR in AI information governance.

Informational Intrusion Risk, Privacy Governance, and Perceived Algorithmic Inequity

IIR captures the users' sense that AI systems are overly monitoring, collecting, or using secondary personal information beyond acceptable limits of privacy and governance (Abraham et al., 2019). The sense of IIR is created when AI systems' data practices violate users' expectations of transparency, control over information collection and use, and fairness in information processing (Hassandoust et al., 2021). From a governance perspective, informational intrusion signals deficiencies in privacy governance mechanisms that regulate how personal data are collected, reused, and protected (Koooper et al., 2011). Drawing on the VAM, IIR is conceptualized as a perceived sacrifice that increases the psychological and governance-related costs of adopting AI applications. As per VAM, when perceived sacrifices exceed perceived benefits, users tend to develop negative value assessments that go beyond usability issues and cover overall assessments of legitimacy, accountability, and fairness (Kim et al., 2007). In the context of AI, perceived intrusive data practices can be considered as a cue that suggests the distribution of informational responsibilities and burdens is not equitable among parties, thus influencing perceptions of algorithmic justice (Majrashi, 2025). Previous studies have found that privacy violations can result in users perceiving algorithms as unfair, especially when some individuals are perceived to be under closer surveillance or face greater risks of data-related issues (Zanotti, 2025). In privacy-conscious domains like the healthcare sector, secondary uses of data for purposes other than the original intended use could be viewed as unfair, particularly when benefits are reaped by a select few but risks are borne by others (Giannouchos et al., 2021). Based on

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

the VAM framework, the factors could influence the perceived privacy sacrifices and perceived unfairness and inequality in AI results (Kim et al., 2007). To influence users' perceptions of unfair risks, benefits, or outcomes that AI systems provide to distinct groups of users, it is hypothesized that IIR influences perceived algorithmic injustice (Hassandoust et al., 2021). When consumers perceive data practices as invasive and inadequately governed, they may conclude that the AI system is biased in favor of some actors, thus increasing perceptions of inequity at the governance level. Therefore, IIR acts as a determinant of perceived risk that utilizes VAM and connects privacy governance inadequacies to fairness and equality judgments in AI systems. When the PAID is high, individuals are likely to form concerns about distributive and procedural justice, which ultimately increases perceptions of inequity at the governance level in AI-enabled environments.

H4: Informational Intrusion Risk has a positive effect on Perceived Algorithmic Inequity and Discomfort.

Consumer Cyber Exposure Risk and Perceived Algorithmic Inequity and Discomfort

CCER is defined as consumers' perceptions of vulnerability to data breaches, hacking, unauthorized access, and insufficient cybersecurity practices in AI systems. In the AI information governance context, cyber exposure risk is an issue when security practices are not being uniformly implemented across platforms and are not being properly communicated, leading to confusion about the protection of personal information and the responsibility for AI system failures. According to the VAM, CCER is a major perceived sacrifice that raises both tangible and intangible costs of adopting AI systems (Kim et al., 2007). VAM argues that when perceived sacrifices of security risks, uncertainty, and potential damage are greater than perceived benefits, users form negative value assessments that generalize to overall judgments of fairness, legitimacy, and equity (Fernando and Dawson, 2009). In AI contexts, insufficient or uneven cybersecurity provisions may signal that system priorities favor operational efficiency over consumer protection, thereby increasing perceptions of governance imbalance. Prior research indicates that asymmetrical cybersecurity safeguards and access control failures can impose disproportionate burdens on certain users or groups, particularly in highly data-intensive and sensitive domains such as healthcare (Hassidim et al., 2017). Such imbalances in risk exposure are some of the factors that contribute to the perception that AI systems are not equitable in their distribution of security risks to users (Gwebu et al., 2020). In addition, the various waves of digitalization, which have been accelerated by crises, have raised consumer awareness of inequitable exposure to cyber risks. From a VAM cost-benefit perspective, elevated Consumer CCER reflects a value imbalance in which governance-related sacrifices outweigh anticipated benefits (Kim et al., 2007). If users feel that they are not equally exposed to cyber risks and that the protection mechanisms are not equal, they may feel that the AI system provides unequal protection to individuals or groups. Therefore, CCER acts as a VAM-based perceived risk factor that leads to PAID, which is an outcome of risk exposure at the governance level.

H5: Consumer Cyber Exposure Risk has a positive effect on Perceived Algorithmic Inequity and Discomfort.

Erosion of User Agency, System Security, and Perceived Algorithmic Inequity

EUA is the degradation or loss of user engagement or involvement in data processing and AI system decision-making (Majrashi, 2025). (Hassidim et al., 2017) stated that the use of safety or security software as control software without providing contestability, transparency, and accountability to users leads to the degradation of user agency in the context of AI information governance (AIG). The VAM declared that the sacrifice of user agency is significant and increases the psychological and governance costs of using an AI system (Kim et al., 2007). VAM argues that when perceived sacrifices in the form of loss of autonomy or lack of contestability of decisions exceed perceived benefits, users form negative value beliefs that also include concerns about fairness and equity (Almulihi et al., 2022). In AI-related applications, security systems that focus on deterrence and surveillance but do not have transparency and end-user involvement in governance processes can lead to decreased perceived value because of user marginalization (Fernando and Dawson, 2009). Previous studies have shown that decreased transparency in monitoring and accountability procedures can perpetuate decreased perceptions of marginalization and inequity (Chernyshev et al., 2019). In critical application areas such as healthcare, decreased transparency in audit trails, centralized security management, and accountability procedures may increase perceptions of inequity, especially when users are at risk without end-user involvement in decision-making procedures (Fernando and Dawson, 2009). Moreover, with the increasing autonomy of AI systems, complex security systems may be invisible to end-users, further reducing their ability to comprehend, question, or impact

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

system-level decision-making (Potluri, 2025). In a VAM cost-benefit analysis, the erosion of user agency is more likely to tip the balance of value toward perceived costs, as the sacrifices associated with governance issues become more significant than the potential benefits (Kim et al., 2007). As users begin to view themselves as passive objects of system decisions, they are likely to perceive AI systems as exercising uneven control and protection in favor of institutional stakeholders rather than individual stakeholders (Chernyshev et al., 2019). As a result, the EUA becomes a VAM-based perceived risk factor that amplifies PAID as a governance-level consequence of reduced user empowerment (Pool et al., 2024).

H6: Erosion of User Agency has a positive effect on Perceived Algorithmic Inequity and Discomfort.

Algorithmic Opacity Risk and Perceived Algorithmic Inequity and Discomfort

AOR is defined as users' perceived inability to understand, evaluate, or assess AI system data processing, decision-making logic, or accountability structures commonly described as the "black box" nature of AI systems. Within AI application information governance, AOR represents a significant risk factor that constrains users' ability to evaluate system fairness, accountability, and legitimacy (Fernando and Dawson, 2009). From a VAM perspective, AOR functions as an interpretive mechanism through which accumulated perceived sacrifices related to AI systems are translated into value-based judgments. The VAM framework posits that when perceived sacrifices outweigh perceived benefits, lack of transparency in AI design, overstated assertions about the capabilities of AI, and a lack of clear explanations about decision logic further cloud users' capacity to evaluate accountability structures, thus exacerbating feelings of inequity. This problem is particularly urgent in high-risk sectors like the healthcare sector, where the AI systems could influence high-risk decisions while the system architecture is opaque (Krokowski and Hirsch-Kreinsen, 2025). The users could therefore feel that there is an imbalance in the allocation of risks and benefits in the AI decision structures (Chernyshev et al., 2019). When perceived sacrifices outweigh perceived benefits, users are more likely to rely on inference rather than evidence-based evaluation in forming their assessments of a technology (Kim et al., 2007). In AI, the capacity of users to assess the decision logic, fairness, unbiasedness, and equity is impeded by the opacity of algorithms. Previous studies have shown that the difficulty of assessing the decision logic and accountability mechanisms of AI is a factor in the perceived unfairness of treatment, especially when users are not aware of the possible biases and discriminatory patterns incorporated in AI systems (Chi et al., 2017) (Choi and Johnson, 2021). A lack of transparency in AI design, overstated claims about the capabilities of AI, and a lack of clear explanation of decision logic further clouds the user's ability to evaluate accountability structures, thus exacerbating feelings of inequity (D'Arcy et al., 2009) (Dutta et al., 2022). This is especially true in high-risk areas such as healthcare, where AI systems could impact high-stakes decisions but the underlying architecture of the system and accountability mechanisms are opaque (Krokowski and Hirsch-Kreinsen, 2025). Users could therefore feel that there is an inequitable allocation of risk and benefit encapsulated in AI decision logic. Across the AI lifecycle, insufficient transparency complicates the evaluation of bias, organizational liability, and corrective mechanisms, thereby increasing uncertainty regarding justice and equity (Odilla, 2024). In highly autonomous and large-scale AI systems, opacity increases users' dependence on assumptions over verifiable evidence, thus increasing the perception of inequity even if the discriminatory outcomes are not directly observable (Fernando and Dawson, 2009). Thus, in the context of AI information governance, AOR is a direct antecedent of PAID, as users feel a lack of transparency and accountability as signs of imbalance in information governance (Recker et al., 2019).

H7: AOR has a positive effect on Perceived Algorithmic Inequity and Discomfort.

User Agency Erosion, Algorithmic Opacity Risk, and Perceived Algorithmic Inequity

In AI information governance, system risk appraisal refers to the identification and evaluation of threats generated by AI systems throughout their lifecycle, considering both their likelihood and potential impact. However, effective governance extends beyond technical safeguards and formal controls; it also requires that users are able to understand, monitor, contest, and meaningfully engage with how AI systems function. When users are faced with limited maneuverability, limited oversight, and limited opportunities to influence AI-driven decisions, their capacity to observe, interpret, and assess system risk is significantly diminished. Based on the VAM, loss of user agency represents a basic perceived sacrifice that contributes to higher psychological and governance costs of AI adoption (Kim et al., 2007). VAM argues that when

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

people feel their efforts are greater than the benefits they will receive, they tend to trust their intuition more than facts when judging system behavior (Choi and Johnson, 2021). In the AI realm, the lack of user agency leads to less engagement with risk identification, monitoring, and management, as well as less ability to hold systems accountable (Paulus et al., 2024). In such a scenario, users are more likely to be exposed to high levels of AOR, which is defined by the perception that AI systems are irrefutable and opaque (Joia and Correia, 2018). Empirical studies show that when user engagement with accountability is low, there is a positive correlation with unpredictability and a lack of participation in governance (Balkan and Akyuz, 2025). While institutions have implemented measures such as cybersecurity ratings, compliance certification, and risk classification to improve system assessment, these measures have been found to be less explanatory and ineffective in overcoming users' perceptions of opacity (Novelli et al., 2024). Moreover, the overstated stories about the autonomy of AI systems could also perpetuate the replacement of human judgment in risk assessment, thus increasing the perception of opacity (Krokowski and Hirsch-Kreinsen, 2025). Studies on team security approaches and distributed decision-making in crisis environments have found that the application of diffuse accountability structures, although enhancing resilience, may also contribute to perceptions of opacity (Johnston et al., 2019). In critical fields such as healthcare and mobile health (mHealth), the presence of complex technical, organizational, and legal protections may make AI outputs and explanations opaque (Harvey and Harvey, 2014). Moreover, the lack of user expertise and inadequate explanation facilities can further contribute to the perception of opacity (Mueller, 2025). In line with the VAM cost-benefit framework, the perception of reduced agency in AI-based decision-making can enhance perceived sacrifices in comparison to the perceived benefits (Kim et al., 2007). Consequently, the user will be more likely to view the lack of transparency as a symptom of an imbalance in governance and inequity. Therefore, the AOR factor will strengthen the PAID factor, which is a governance-level outcome in which opacity is a mediator between reduced agency and perceptions of injustice in AI systems (Lämmermann et al., 2024).

H8: AOR mediates the relationship between Erosion of User Agency and Perceived Algorithmic Inequity and Discomfort.

Conceptual Framework

The proposed Figure 1 illustrates consumer perception of AI application information governance by explaining how different perceived risks shape negative algorithmic outcomes. Specifically, IIR, CCER, and EUA act as key antecedents influencing consumers' perceptions of algorithmic AOR and PAID. Paths H1-H3 indicate that when consumers perceive excessive data collection, heightened cyber vulnerability, and reduced control over AI-driven decisions, they are more likely to view AI algorithms as opaque and non-transparent. This opacity, in turn, directly increases PAID (H7), suggesting that a lack of explainability and visibility in AI systems fosters feelings of unfairness, bias, and unease. Additionally, paths H4-H6 demonstrate that these risk perceptions also exert direct effects on PAID, highlighting that consumers' concerns are not solely explained by opacity but also stem from broader governance and autonomy issues. The model further proposes mediation effects, where AOR partially or fully mediates the relationship between governance-related risks and negative consumer outcomes (H8), underscoring the central role of transparency in AI information governance. Overall, the framework emphasizes that weak information governance in AI applications amplifies consumer discomfort both directly and indirectly through perceived algorithmic opacity.

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

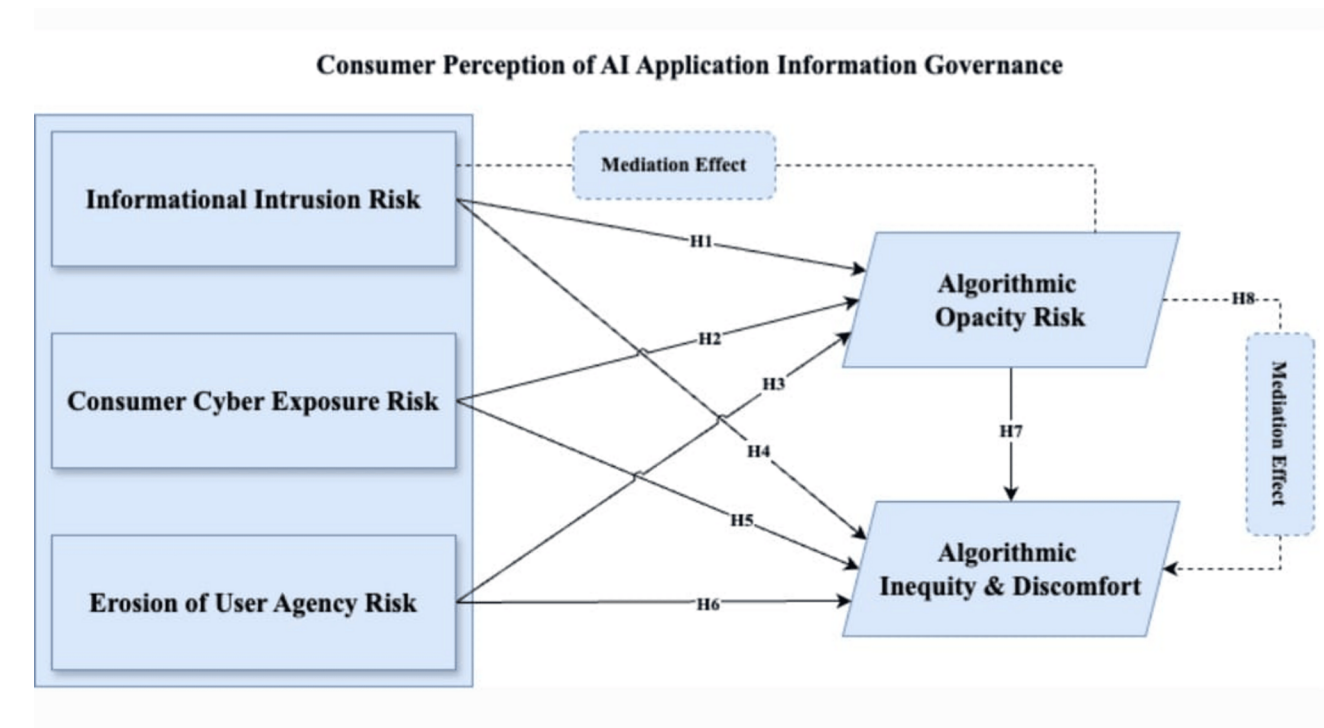


FIGURE 1: Conceptual model

Source: Authors

Research Method

Data collection

Data collection (Figure 2) was conducted through social media channels, namely WhatsApp and Facebook. These platforms host professional communities, such as the Pakistan Students and Scholars Association South Korea (PSAK), with approximately 8,000 followers, and SEEKERS, a network of educated students and professionals, with approximately 17,000 followers on both platforms. We used the questionnaire, and a total of 300 responses were collected for data analysis. SPSS and AMOS software were used. This figure shows a line chart titled "Survey Responses Over Time," depicting daily response counts from early November to early December 2025. Dates are listed along the horizontal axis, and the number of responses is shown on the vertical axis. In the beginning stages, the data points sit low in the gutters, with just a few totals represented from 1 to 6 and an occasional odd day represented by the number 13. As the month of November progresses, the data points jump upward in fits and starts, showing a fast climb upward and a dip downward. The high points are reached toward the end of the month, with data points brushing against 30. Next comes a dip downward, followed by upward surges and a climb to the start of the month, with the data points pushing to 20 to 28 by the end of this sequence of events.

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

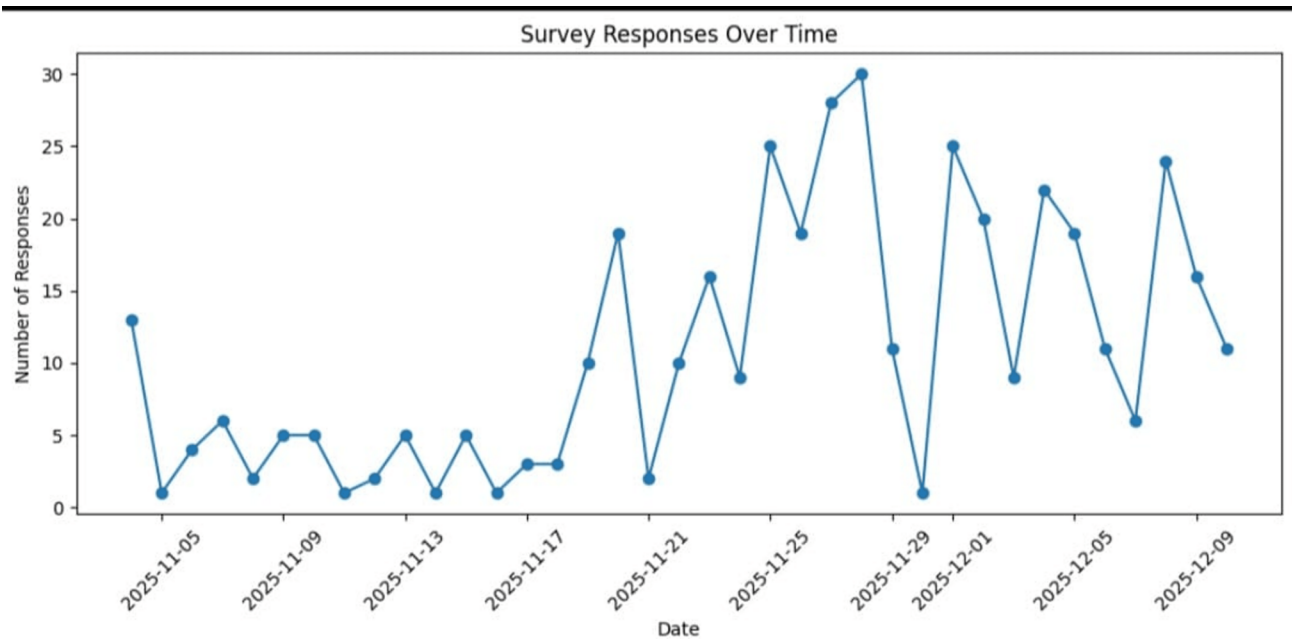


FIGURE 2: Data collection timeline

Source: Authors

Measurement items

Table 7 shows the data collection purposes; the structured questionnaire was used to measure the operationalization of the proposed constructs in the research model. The questions in the structured questionnaire are based on the scales used in the past, and the questions are tailored to the context and used in the context of information governance for the application of AI. The details of the measurement items, definitions, and sources are listed in Table 7. The five-point Likert scale is used for data collection, and the scale ranges from 1 to 5, which stands for "strongly disagree," "disagree," "neither," "agree," and "strongly agree," respectively. The respondents are given the opportunity to specify the degree of their perception within this context and with respect to AI risks. The IIR is monitored with respect to four indicators, namely IIR1, IIR2, IIR3, and IIR4. The indicators measure how people feel about information collection, information usage outside of original intent or permission, information ownership confusion, and information retention. CCER has another set of four indicators: CCER1 refers to consumer concerns about hacking; CCER2 refers to consumer concerns about information usage; CCER3 refers to consumer concerns about cybersecurity; and CCER4 refers to consumer concerns about inequality of cyber harm. The EUA has another set of four indicators or dimensions: EUA1 to EUA4. The EU indicators measure consumer perception of losing control in AI-related decisions, reliance on decisions taken autonomously by AI, lack of disconfirm ability of AI-related information, and lack of explanation about information usage. It is characterized by four measures: AOR (AOR1-AOR4), which refers to the perception of the consumer that he or she does not understand the AI-related decision, especially when there is the occurrence of cyber vulnerability, or access issues, etc. Lastly, perceived algorithmic inequity and discomfort have another set of four indicators: PAID1-PAID4, which refers to issues of fairness, enhancing bias, suspicion of unintended effects of AI-related information usage, and erosion of trust as a result of the mysterious nature of AI models.

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

Measurements Items	Definitions	Questions	Sources
Informational Intrusion Risk	The concern that AI systems may collect and utilize too much personal data or do so inappropriately.	IIR 1—If their personal information is being collected by applications based on AI in excessive amounts, I do not think their algorithm is clear to me. IIR 2—If my personal information is used by the applications related to artificial intelligence in a way that I did not originally agree to, then I am having difficulty understanding the way by which the applications' decisions are made. IIR 3—I am unsure who can access my personal data in AI applications, and this makes them seem non-transparent and opaque to me. IIR 4—The longer AI applications store my personal data, the less transparent the data processing and decision-making procedures seem to me.	(Recker et al., 2019)
Consumer Cyber Exposure Risk	The concern about cybersecurity threats that come with AI-enabled applications.	CCER 1—Since artificial intelligence applications wield significant influence over my personal information, I am quite uncomfortable about the treatment of users by the algorithmic decisions. CCER 2—Concerns regarding exposure of my personal data because of hacking and misuse make me feel apprehensive regarding the fairness of decisions made by AI technology. CCER 3—If the cybersecurity breakdowns happen in AI, I am concerned that their algorithmic result might harm me. CCER 4—The risk that some unauthorized individuals may take advantage of AI systems gives	(Weber, 2020)

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

		me the concern that some users may be negatively affected by an inequitable AI-related decision.	
Erosion of User Agency	The feeling of users losing control and influence over the decisions led by AI.	EUA 1—I strongly believe that because of the use of my own data in numerous AI applications, I have no control over the decisions that are made for me. EUA 2—AI applications have the ability to make critical decisions without human interference. The decision process in these applications is opaque to me. EUA 3—When personal data is extensively utilized, challenging decisions made by artificial intelligence algorithms become difficult. EUA 4—The manner in which the AI applications utilize my personal data is not elucidated sufficiently so that I can comprehend their reasoning.	(Giannouchos et al., 2021)
Algorithmic Opacity Risk	The perceived inability to understand how AI algorithms generate decisions.	AOR 1 —In current AI applications, when my data are threatened by cyberattacks, I feel a lack of control over the decisions that could discriminate against users. AOR 2— In current AI applications, when my data are threatened by cyberattacks, it is hard for me to understand decisions that could discriminate against users. AOR 3—Due to the weaknesses that exist in cyberspace and that are out of my control, I am uneasy with the way decisions are made by applications concerning me by artificial intelligence. AOR 4—If AI systems are vulnerable to unauthorized access, I am concerned that the results of these algorithms might not always favor all individuals.	(Hassidim et al., 2017)

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

Perceived Algorithmic Inequity and Discomfort	The perceived unfairness of AI-driven outcomes and the discomfort they cause.	PAID 1—In cases when AI is non-transparent, it may lead to unfairness due to hidden biases. PAID 2—Because AI decision-making is opaque, AI-driven results could lead to an exacerbation of existing social inequality. PAID 3—The black box nature of AI applications makes me suspicious of their potential unintended outcomes. PAID 4—As a result of AI "black boxes," failures decrease my trust in the use of AI technologies.	(Hassandoust et al., 2021)
---	---	--	----------------------------

TABLE 1: Measurement items

Source: Authors

Data analysis tools and techniques

The analysis was carried out through careful scrutiny following the scientific method, utilizing widely recognized statistical application toolsets. The toolset deployed was SPSS Statistics software. Phase 1 included ensuring the accuracy of the information and handling missing data, as well as carrying out a preliminary assessment of the respondents. The methodology deployed in collecting information for analysis consisted of five options in a Likert scale, ranging from 1 for "strongly disagree" to 5 for "strongly agree." The system anticipates collecting the perceptions of the respondents concerning information risk associated with AI implementation. To test the measurement model and assess the hypothesized relations, AMOS was employed to conduct both CFA and SEM, as it can help test construct validity, including convergent and discriminant construct validity, where the relations among latent variables can simultaneously be tested and the relations of mediation can be investigated in terms of how the threat of opaque algorithms can influence the perception of risk in governance and the related sense of inequity and algorithmic discomfort. In general, by using SPSS software together with AMOS 26, a comprehensive analytical setting is created, which confirms the reliability, validity, and rigor of the research findings included in this research.

Results And Discussion

Data analysis and results

Demographic Analysis

Table 2 illustrates the characteristics of consumers. This study included the demographic questions related to gender, age, education, profession, and income to know about the respondents' characteristics. The collected data at the end showed that the respondents (n = 300) indicate a male-dominated sample, with 74.0% male (n = 222) and 26.0% female (n = 78) participants. The gender breakdown implies that the results could better represent views from a male perspective. Moving to age, most respondents fall within the 21-30 age group (55.0%) and then the 31-40 age group (38.0%), meaning that most respondents fall within the young to middle-aged demographic. Fewer respondents fall

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

within the 10-20 age group (3.0%), 41-50 age group (3.7%), and over 50 age group (0.3%), meaning that there is poor representation from the older generation. For educational qualifications, most people fall within the advanced educational level and include those with Ph.D. educational levels (43.0%) or those in pursuit of a Master's (39.0%) level, with 18.0% in pursuit of a Bachelors level. This pattern reveals that the respondents are highly educated, which is suitable for the investigation of complex constructs like algorithmic risk and perceptions. The occupational characteristics of the respondents reveal that almost half of the population is comprised of students (47.7%) followed by those associated with research (25.3%), which reveals good representation from academic as well as research communities. The other occupations include engineers (7.0%), private workers/businessmen (3.0%), artists (1.3%), and others (15.7%), varying somewhat on the basis of their occupation. Concerning monthly income, the largest group of respondents falls between USD 500-1000 (39.7%), followed by those earning between USD 1000 and 1500 (26.3%). The rest, constituting the smaller group, earn between USD 1500 and 2000 (7.7%), followed by a few who earn between USD 4000 and 8000 (1.0%). The "Others" group, constituting 25.3%, indicates variability or nondisclosure among a substantial number of respondents regarding their earnings. The demographic breakdown presents a considerably young, highly educated, and academically engaged group with moderate earnings, appropriate for studying impressions about algorithmic transparency, agency, or lack of fairness. However, generalization remains an important concern, as the sample is less representative of older or less-educated populations, and the results should therefore be interpreted cautiously. The analysis adds emphasis on the fact that, due to the high percentage of men in the sample being highly educated, there is an influence on how AI governance risks might be perceived. This is because men in general tend to be more knowledgeable about technology than women. Moreover, the fact that there is a high percentage of respondents in the younger age groups (21-40) indicates that the findings have been more representative of digitally experienced users. This is beneficial in gaining more insights into AI interactions. However, there is also a high percentage of students in the sample, which indicates that the findings have been more representative of a cognitively engaged population. This is beneficial in assessing complex constructs but might be biased toward more analytical evaluations of risks. According to the authors, the demographics chart provided in the category "Others" (15.7%) includes all categories of people: educated, private sector employees, government (Govt.) sector employees, businessmen, artists, and laborers. Nonetheless, the current research intends to explore the consumers' views on AI-powered systems, focusing on users who are actively involved in this process. Involvement of all categories of people in using AI systems means that their perspectives are essential for further understanding this phenomenon. According to the authors, there is not any issue with external validity. As the current study is a small-scale study, it can be constrained by the sample size. This is the reason it has focused on a small sample size from the population. According to the authors, all the descriptive texts are based on the literature review and research process, such as methodology, data analysis, and overall results and discussion.

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

Demographical Factors	Sample (N = 300)	Percentage %
Gender		
Male	222	74.00%
Female	78	26.00%
Age		
10-20	9	3.00%
21-30	165	55.00%
31-40	114	38.00%
41-50	11	3.70%
50 Above	1	0.30%
Qualification		
Bachelors Student	54	18.00%
Masters Student	117	39.00%
PhD and others	129	43.00%
Profession		
Student	143	47.70%
Researcher	76	25.30%
Engineer	21	7.00%
Private Employee/Businessman	9	3.00%
Artists	4	1.30%
Others	47	15.70%
Income (Monthly)		
500USD-1000USD	119	39.7
1000USD-1500USD	79	26.3
1500USD-2000USD	23	7.7

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

4000USD-8000USD	3	1
Others	76	25.3

TABLE 2: Respondents characteristics

Source: Authors

Descriptive statistics, reliability, and correlations

Table 3 shows the results of the characteristics of the sample, reliability analysis, and relationships of the variables included in the research using the sample of 300 people. In most of the cases, averages are at a mid-to-high level, between 13.58 and 15.73 points. This shows positive results indicating good perception of these constructs. The most significant factor found was IIR, with the top scores of the sample, with an average of 15.73 points and a standard deviation of 3.37. This shows that respondents are highly concerned about the amount of collected information during AI application. In contrast, the weakest perception of respondents and of the whole sample was found for PAID - the least results of these constructs - with an average of 13.58 points and a rather large standard deviation of 4.3, showing wide variation of opinions, mostly negative. All constructs show excellent reliability, with Cronbach Alpha results of 0.921, 0.979, far above the required 0.70. The result of the correlation analysis indicates clear and significant relationships among the variables at a 0.01 confidence interval. IIR is noteworthy in that it has a significant relationship with CCER at $r = 0.882$. This indicates that when there were apprehensions about the intrusive use of user data, apprehensions about cybersecurity exposure were always present. EUA indicates a high relationship with AOR at $r = 0.722$ and PAID at $r = 0.700$. This suggests that when there was a lack of EUA, perceptions of non-transparency and algorithmic unfairness always accompanied AOR. However, AOR and PAID are also strongly related at $r = 0.851$. This suggests that perceptions of high opacity also strongly corresponded with perceptions of unfairness in the algorithms. This suggests strong reliability in terms of relationships within the constructs and strong empirical support from a conceptual framework.

Variable	Mean	SD	IIR	CCER	EUA	AOR	PAID
IIR	15.73	3.37	<i>-0.938</i>				
CCER	15.28	3.32	.882**	<i>-0.921</i>			
EUA	14.74	4.11	.452**	.548**	<i>-0.979</i>		
AOR	14.34	4.18	.406**	.443**	.722**	<i>-0.966</i>	
PAID	13.58	4.3	.395**	.431**	.700**	.851**	<i>-0.963</i>

TABLE 3: Descriptives, reliability, and correlations

Source: Authors. Note: $n = 300$; $p < 0.01$. Cronbach's alpha = Italic values. SD, Standard Deviation; IIR, Informational Intrusion Risk; CCER, Consumer Cyber Exposure Risk; EUA, Erosion of User Agency; AOR, Algorithmic Opacity Risk; PAID, Perceived Algorithmic Inequity and Discomfort

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

Confirmatory factor analysis model structure

As shown in Table 4, the measurement model above indicates that the reliability of this model is strong and the convergent validity is good, while the overall model fits well. All standardized factor loadings range from 0.766 to 0.971, well above the cut off of 0.70. Thus, it is evident that these measures are reliable in measuring their respective constructs well. As for reliability, Cronbach's Alpha is also well above the cut off of 0.70, ranging from 0.857 to 0.959. Moreover, Composite Reliability ranges from 0.916 to 0.979, confirming well-developed reliability among these or across all of the constructs. As for convergent validity, the range of AVEs from 0.734 to 0.919 is well above the cut off of 0.50. Besides, with regards to model fits, the measurement model has shown adequate overall model fit. As per Hair et al., and with data shown in model fits above, the measurement model's parsimonious fits are well within the range of 3 to 5 as χ^2/df is equal to 2.276; with GFI and AGFI ranging from 0.901 to 0.865, it is well above the cutoff point of 0.80; and incremental fits range from 0.960 to 0.977, surpassing well the cut off of 0.90. Besides, with RMSEA and RMR of 0.065 and 0.040, respectively, they are relatively within the cut off of 0. The high standardized factor loadings of all constructs confirm that all indicators are significant representatives of the respective underlying theoretical construct, ensuring that the measurement of IIR, CCER, EUA, AOR, and PAID is grounded in theory and consistent with the proposed framework. Additionally, the high values of Cronbach's Alpha and Composite Reliability for all constructs confirm that all indicators are coherent representatives of the respective underlying dimension of AI governance risk. Moreover, the satisfactory values of AVE also confirm convergent validity, which supports the assumption that each construct represents a unifying concept while still being statistically different from other constructs. The distinction between similar constructs is also crucial in the case of EUA and AOR, as there is theoretical similarity but still statistical distinction. Overall, these results not only go beyond statistical results but also confirm that the measurement model is in line with theoretical expectations, thus creating a solid and reliable foundation to examine the structural relationships and validate the proposed model in the context of AI governance risks.

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

Constructs	Indicators	Std. Factor Loadings	Purpose of Deploying the Construct	Cronbach's Alpha (CA)	Composite Reliability (CR)	AVE
Informational Intrusion Risk (IIR)	IIR 1	0.828	We used IIR to understand whether users feel AI applications collect too much personal information.	0.922	0.93	0.782
	IIR 2	0.92	We used IIR to examine concerns about personal data being used without original consent.			
	IIR 3	0.919	We used IIR to assess whether uncertainty about data access and storage makes AI algorithms seem unclear.			
	IIR 4	0.882	We used IIR to explain how privacy intrusion increases perceptions of non-transparency in AI systems.			
Consumer Cyber Exposure Risk (CCER)	CCER 5	0.914	We used CCER to measure users' fear of hacking, data misuse, and unauthorized access.	0.828	0.916	0.73

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

	CCER 6	0.922	We used CCER to examine whether cybersecurity threats affect fairness perceptions of AI decisions.			
	CCER 7	0.766	We used CCER to understand how weak cybersecurity makes users feel exposed and unprotected.			
	CCER 8	0.83	We used CCER to link cyber risks with discomfort and concern about harmful AI outcomes.			
Erosion of User Agency (EUA)	EUA 9	0.92	We used EUA to assess whether users feel a lack of control over AI-driven decisions.	0.82	0.976	0.919
	EUA 10	0.972	We used EUA to examine concerns about AI making important decisions without human involvement.			
	EUA 11	0.966	We used EUA to understand difficulties users face in challenging AI decisions.			

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

	EUA 12	0.922	We used EUA to explain how reduced user control makes AI systems feel opaque and unaccountable.			
Algorithmic Opacity Risk (AOR)	AOR 13	0.965	We used AOR to measure how AI decision-making feels unclear or like a "black box."	0.925	0.919	0.852
	AOR 14	0.966	We used AOR to understand why users struggle to see how AI decisions are made.			
	AOR 15	0.889	We used AOR to examine whether lack of transparency increases fairness and trust concerns.			
	AOR 16	0.878	We used AOR to test opacity as a link between loss of user agency and perceived unfairness.			
Perceived Algorithmic Inequity and Discomfort (PAID)	PAID 17	0.926	We used PAID to capture perceptions that AI decisions may be unfair or discriminatory.	0.915	0.951	0.828
	PAID 18	0.925	We used PAID to measure users'			

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

			uneasiness and discomfort with AI decision-making.			
	PAID 19	0.918	We used PAID to assess concerns that AI outcomes may not benefit all users equally.			
	PAID 20	0.826	We used PAID to evaluate how opacity and cyber risks reduce trust in AI technologies.			

TABLE 4: Confirmatory factor analysis

Source: Authors

Assessment of structural model fit

As shown in Table 5, the findings show that the proposed model provides a good and strong fit for the data. Even though the chi-square statistic is statistically significant ($\chi^2 = 352.720$, $df = 156$, $p < .001$) this is not uncommon in SEM modeling. This is especially true in situations where the data are sampled from a moderate- to large-sized sample. This is due to the high sensitivity of the chi-square statistic to large sample size and slight model misspecification. Therefore, the focus is given to alternative indices. The normed chi-square ($\chi^2/df = 2.276$) is less than the established limit of 3.0. This indicates acceptable levels of parsimony and fit. In relation to absolute fit, the GFI index has a value of .906, higher than the cut-off of .90, while the AGFI has a value of .865, significantly higher than .80, thus demonstrating that the model presents sufficient fit to the observed covariance matrix. With respect to incremental (comparative) fit, all of these NFI (.960), IFI (.977), TLI (.972), and CFI (.977) are well above the critical level of .95, thereby indicating that the model shows excellent improvement over the null model. In relation to the RMSEA, its index of .065, although lower than .08, is also near .06, suggesting that the fit is good in approximating the population covariance matrix. Hence, these results demonstrate that the model has achieved sufficient fit, suggesting that it can be used to test the hypotheses.

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

Fit Index	Obtained Value	Interpretation
Chi-Square (χ^2)	352.720	Sensitive to sample size
Degree of Freedom	156	Model complexity indicator
χ^2/df (CMIN/df)	2.276	Acceptable fit
P-value	.0001	Significant (common in large samples)
GFI	.906	Acceptable fit
AGFI	.865	Acceptable fit
RMR	.040	Good fit
NFI	.960	Good fit
IFI	.977	Excellent fit
TLI	.972	Excellent fit
CFI	.977	Excellent fit
RMSEA	.065	Acceptable

TABLE 5: Model Fit Indices

Source: Authors. GFI, Goodness-of-Fit Index; AGFI, Adjusted Goodness-of-Fit Index; RMR, Root Mean Square Residual; NFI, Normed Fit Index; IFI, Incremental Fit Index; TLI, Tucker–Lewis Index; CFI, Comparative Fit Index; RMSEA, Root Mean Square Error of Approximation

Discriminant validity

From Table 6, it is clear that discriminant validity exists, together with the relationship between the different variables in this study. The diagonal scores, ranging from 0.857 to 0.959, represent the square root of the total square root of each of the constructs' AVE score, which is situated clearly above the off-diagonal scores. This verifies that the study's measures of risk types, i.e., AOR, IIR, CCER, EUA, and PAID, are clear distinct from one another. Furthermore, the study found the following connections to be consistent with theoretical predictions: "AOR" had strong theoretical relationships with EUA (0.762) and with PAID (0.906), which explains that the less control the consumer has, the more uncomfortable the AI system will be perceived to be. IIR and CCER had moderate-to-strong positive relationships with the study's measures of Algorithmic Opacity (0.459 and 0.469), revealing the connection to the theory that general concerns over AI system privacy drive opinions on the lack of transparency in the system's associated algorithms.

How to cite this article:

	Algorithmic Opacity Risk	Informational Intrusion Risk	Consumer Cyber Exposure Risk	Erosion of User Agency	Algorithmic Inequity and Discomfort
Algorithmic Opacity Risk	0.924				
Informational Intrusion Risk	0.459	0.884			
Consumer Cyber Exposure Risk	0.469	0.954	0.857		
Erosion of User Agency	0.762	0.468	0.546	0.959	
Algorithmic Inequity and Discomfort	0.906	0.439	0.443	0.739	0.915

TABLE 6: Discriminant validity

Source: Authors

Path coefficient estimates and hypothesis testing results

The results in Figure 3 show that each construct measures up well, and the proposed associations indeed occur in the real world. A large relationship exists in the finding that IIR and AOR go together ($\beta = 0.88$), indicating that if a consumer feels that the data collection was excessive, the algorithm was non-transparent. In addition, if the perception of consumer cyber risks perception exists ($\beta = 0.50$), and the erosion of control perception exists ($\beta = 0.55$), it demonstrates that consumer risks contribute to opacity perception. The direct effects on perceived inequity in algorithms and levels of unease show positive, though limited, effects in terms of risks of informational intrusion ($\beta = 0.16$), cyber exposure ($\beta = 0.17$), and losing control ($\beta = 0.15$), while the effect of algorithmic opacity is strong in terms of algorithmic inequity and discomfort ($\beta = 0.47$), as it significantly implies the degree to which the presence of AI is perceived to promote inequity and discomfort if it is opaque to the social subjects. Based on these outcomes, it can be concluded that the mediating role of the risk of opacity in algorithms in the relationship between perceived risks by consumers and inequity involves transparency as a cornerstone in the governance of information in AI systems.

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

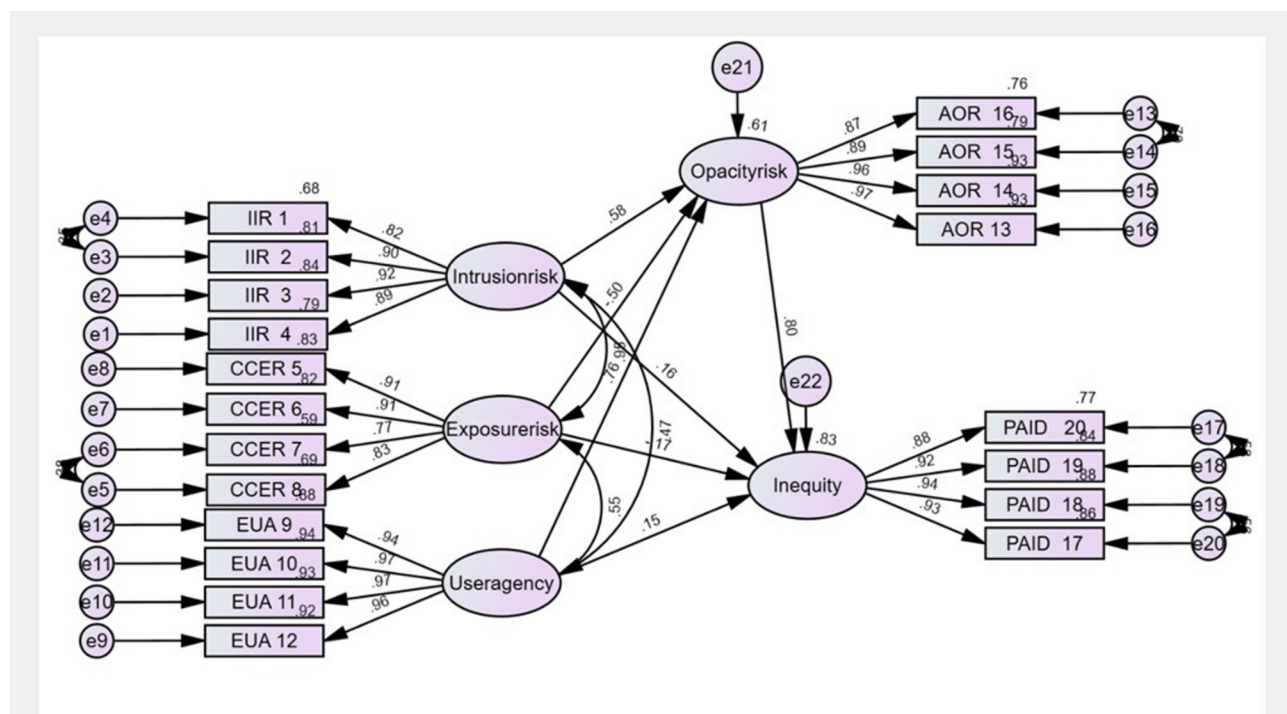


FIGURE 3: Structural equation model

Source: Authors. IIR, Informational Intrusion Risk; CCER, Consumer Cyber Exposure Risk; EUA, Erosion of User Agency; AOR, Algorithmic Opacity Risk; PAID, Perceived Algorithmic Inequity and Discomfort

Hypothesis testing

The results of the hypothesis test show a high level of significance of the research model that includes seven proposed relationships (Table 7). All hypotheses reveal that the direction of the proposed relationship is positive and statistically significant for all paths ($p = 0.001$). Thus, we can say that all suggested relationships between AI risk, algorithmic opacity, and consumers' perception of inequity in algorithms are reliable and significant. It is important to note that the standardized path coefficient (β) of each path is between 0.211 and 0.279. Hence, we may conclude that all the variables positively contribute to the explanation of consumers' perceptions of AI systems. Therefore, we can say that the concerns of consumers about privacy, cybersecurity, and loss of control are not separate factors; they are interrelated and impact the overall evaluation of the AI equity and transparency. With respect to the VAM, risk perceptions may be considered as the perception of the sacrifices that are made by users of AI technology. Thus, the increased concern about the risk of using AI applications increases the concerns about their work. The first hypothesis (H1) was formulated with regard to the relationship between IIR and AOR. According to the obtained results, there is a significant positive relationship ($\beta = 0.221$, $p = 0.001$), which means that high perception of informational intrusion correlates with high perceptions of algorithmic opacity. As a result, it becomes obvious that people get more worried about AI transparency when they perceive AI systems as those that gather too much personal information or monitor their online actions. Nowadays, many AI systems are designed to collect vast amounts of user information to make recommendations, predictions, and automated decisions for the user. The results confirm the existence of a positive impact of the variable on algorithmic inequity ($\beta = 0.223$, $p = 0.001$), meaning that when people worry about their privacy, they think that the system is acting unfairly to them. There is a probability that consumers feel that excessive gathering and use of their personal information leads to the discrimination of users. For instance, when the information that is used by an AI system is incomplete, wrong, and biased, there is a chance of making decisions that affect particular individuals negatively. The fifth hypothesis (H5) looked at the influence of consumer cyber exposure risk on algorithmic inequity. Examples of how perceived risk impacts the perception of AI based on the transparency and fairness aspects. In practical terms, AI developers, companies, and

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

governments should focus on developing explainable AI, responsible data management, cybersecurity, and empowering consumers to make informed decisions. As it can be seen from the results, there is a significant positive association ($\beta = 0.226$, $p = 0.001$), which means that the risk of cybersecurity influences the perception of potential unfairness of AI systems. It is likely that consumers worried about the cyber risks will wonder if AI systems are able to provide adequate protection of consumer data and if decisions made by algorithms are based on secure data. In this regard, the importance of providing cybersecurity protection becomes not only a technical aspect but also a way of promoting perceptions of the absence of unfairness. The sixth hypothesis (H6) examined the relationship between EUA and algorithmic inequity. The findings confirm a significant positive relationship ($\beta = 0.211$, $p = 0.001$). This implies that when users think that they have little control over decisions made by AI, they view such decisions as being unfair. The expectation among users is that the technology will be used in a manner that will complement their decisions and not replace the judgment of the users. Decisions made by the system without any input from the user, without the provision for explanations, and without providing an opportunity to appeal such decisions, make users feel marginalized during the decision-making process.

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

Hypothesis No.	Hypothesized Relationship	Path	β	p-value	Result
H1	Informational Intrusion Risk positively influences Algorithmic Opacity Risk	Intrusion Risk → Opacity Risk	0.222	0.001	Supported
H2	Consumer Cyber Exposure Risk significantly influences algorithmic opacity risk.	Exposure Risk → Opacity Risk	0.228	0.001	Supported
H3	Erosion of User Agency positively influences Algorithmic Opacity Risk	User Agency → Opacity Risk	0.221	0.001	Supported
H4	Informational Intrusion Risk positively influences Algorithmic Inequity	Intrusion Risk → Inequity	0.223	0.001	Supported
H5	Consumer Cyber Exposure Risk significantly influences algorithmic inequity.	Exposure Risk → Inequity	0.226	0.001	Supported
H6	Erosion of User Agency positively influences Algorithmic Inequity	User Agency → Inequity	0.211	0.001	Supported

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

H7	Algorithmic Opacity Risk positively influences algorithmic inequity.	Opacity Risk → Inequity	0.278	0.001	Supported
----	---	----------------------------	-------	-------	-----------

TABLE 7: Hypothesis analysis and results

Source: Authors

The seventh hypothesis (H7) was designed to investigate the relationship between AOR and algorithmic inequity. The findings showed that there was the highest relationship among all hypotheses tested ($\beta = 0.279$, $p = 0.001$). It is evident that transparency is a significant factor for determining AI fairness. The inability to provide explanations for users to comprehend how the algorithms work causes decisions to be made that make possible discrimination and inequitable decision-making, and the lack of explainability makes it impossible to verify whether the data applied by the algorithm is reliable, hence increasing uncertainties. It means that it is vital to enhance transparency and explainability to reduce inequities. The structural model results demonstrate that privacy concerns, cybersecurity risks, and reduced user autonomy are important predictors of consumers' perceptions of AI opacity and inequity. Among all factors, algorithmic opacity has the strongest influence on perceived inequity, indicating that transparency is a key mechanism connecting AI decisions to fairness concerns. This highlights the need to take into consideration the comprehensibility, controllability, security, and fairness aspects of AI, apart from how functional and valuable the AI is to consumers. With regard to the theoretical aspects, the current research builds on the VAM, providing insights that can help reduce concerns about AI.

Indirect mediation effect

The aim of the mediation analysis is to test the role of AOR as the mechanism underlying the connection between EUA and PAID (Table 8). In such a model, EUA serves as the independent variable, AOR is a mediating variable, while PAID acts as the dependent variable. In general, the mediation analysis is needed to identify whether the impact of the feeling of a loss of control over AI on the perceived algorithmic inequity and discomfort occurs partly via the perception of opacity of AI. Indirect effect from the analysis is equal to $\beta = 0.228$. Indirect effect can be understood as the joint effects of EUA on AOR and AOR on PAID. This effect is usually presented as the product of both paths $a \times b$. Positive value of 0.228 means that high levels of EUA will cause high levels of PAID due to AOR. This fact shows that the more users lack control over AI usage, the greater their perceptions of algorithmic opacity will be. The statistical significance of the indirect effect is proved by the fact that its p-value is below 0.001. This fact proves that the statistically significant indirect effect does not equal zero. The reported 95% confidence interval equals 0.11, 0.12. The fact that both limits of the interval are positive and the interval does not include zero means that the indirect effect is statistically significant according to the standard bootstrap mediation criterion. Namely, the interpretation of the above findings may follow a three-step pattern. First, EUA is a sense of users' inability to control or have any influence on AI systems. Once users feel their inability to control something, they start worrying about the functioning of AI systems and their decisions. Secondly, such concerns might increase the risk of Algorithmic Opacity, meaning that users perceive the processes, decision-making, and reasoning of the AI systems as incomprehensible and obscure. Finally, AOR will result in higher feelings of algorithmic inequality and discomfort when AI decisions are not comprehensible, accessible, and challengeable for users. In other words, AOR is the explanatory factor that exists between EUA and PAID. These findings are not only about the association between EUA and PAID but also point to the fact that this association exists through the perceptions of algorithmic opacity held by users. This finding is theoretically significant in the sense that the absence of user control contributes to these negative perceptions partially due to algorithmic opacity perceptions. For instance, when an individual uses an AI-driven recruitment system while having limited capability to interpret and affect the process of evaluation of the

How to cite this article:

application by the system, it is likely for such an individual to have decreased perceptions of control. When such a system does not provide sufficient information on how the decisions are made, an individual is likely to feel that the algorithm lacks transparency. In turn, lack of transparency may cause an increased perception of unfairness or discrimination. The result marked “Mediation supported” is therefore consistent with the statistical evidence provided in the table. The significant positive indirect effect, $p < 0.001$, and confidence interval that does not cross zero collectively support the conclusion that AOR plays a significant mediating role in the relationship between EUA and PAID. However, it is necessary to differentiate between mediation and causation. Statistical significance of mediation means that there exists an relationship between the variables along the proposed indirect path; however, the cross-sectional survey study does not provide the possibility to conclude that the EUA causes the AOR or vice versa. Mediation analysis sought to establish whether AOR would mediate the relationship between EUA and PAID. The results revealed a statistically significant and positive indirect effect ($\beta = 0.228$, $p < 0.001$, 95% CI [0.11, 0.12]). Because the confidence interval does not contain zero, there is a statistically significant indirect effect, implying that there is a mediation relationship. The results imply that there is a correlation between EUA and AOR, which, in turn, correlates with PAID. However, given the cross-sectional nature of the study, the findings should be interpreted as evidence of an indirect association rather than definitive causal mediation.

Mediation Hypothesis Effect Type	Indirect Path	β (Indirect Effect)	95% CI - LL	95% CI - UL	p-value	Result Status
Indirect Effect (a $\times$ b)	Erosion of User Agency $\rightarrow$ Algorithmic Opacity Risk $\rightarrow$ Algorithmic Inequity and Discomfort	0.228	0.11	0.12	< 0.001	Mediation Supported

TABLE 8: Mediation analysis result
Source: Authors. CI, Confidence Interval; LL, Lower Limit; UL, Upper Limit

Discussion

The mediation analysis was carried out in order to investigate the mediating role of AOR in the influence of EUA on PAID. The purpose of the mediation analysis is to find out whether the effect of decreased user agency (i.e., decrease in the user’s control, influence, and freedom of initiative regarding the AI system) on the perceived inequity and discomfort in relation to algorithms is mediated, at least partially, by the user’s perception of algorithms’ opacity. The results indicate that there is strong indirect effect (a $\times$ b), which is $\beta = 0.228$. The statistically significant indirect effect shows that the erosion of user agency has a meaningful impact on PAID through AOR. Therefore, when users feel that they have limited control or influence on AI systems, this feeling has an impact on their perception of the opacity of algorithms. When users have increased concerns about the opacity of algorithms, this leads to greater perceptions of algorithmic inequity and discomfort. The significance of the indirect pathway is particularly important because mediation analysis goes beyond examining whether two variables are simply related. It helps explain how or through what mechanism one variable influences another. In the present study, EUA represents the perceived loss of user control or autonomy, AOR represents concerns regarding the lack of transparency and understandability of AI decision-making, and PAID represents users’ perceptions of unfairness, inequity, and discomfort associated with algorithmic systems. The mediation results therefore

suggest a sequential relationship in which reduced user agency is associated with greater perceptions of algorithmic opacity, which subsequently contributes to perceived algorithmic inequity and discomfort. The $\beta = 0.228$ indirect effect indicates a positive mediation relationship. This means that increases in EUA are associated with increases in PAID through AOR. The positive direction is theoretically meaningful because users who feel that they cannot adequately influence or understand AI-supported processes may become more concerned about the way algorithms operate. When the decision-making process is unclear, users may have difficulty determining whether an AI system is treating them fairly. Consequently, uncertainty about the algorithm can translate into perceptions of inequity and discomfort. With a value less than 0.001, the p-value strongly indicates statistical significance of the indirect effect. As the p-value is well below the significance level of 0.05, it allows the rejection of the null hypothesis that there is no indirect mediation effect. Therefore, the mediation effect found in Table 8 is statistically significant. The finding confirms the mediation hypothesis stated and indicates that AOR is more than an associated variable but a significant mediator linking EUA to PAID. The confidence interval in the table further confirms the mediation finding. For the 95% confidence interval, it ranges from 0.11 to 0.12. Since both the lower and upper limits of the confidence interval are greater than zero and the value of zero is not contained in the confidence interval range, it means that there is a significant difference between the indirect effect and zero.

This provides more evidence regarding the significance of the mediation effect. In mediation analysis, the confidence interval, which does not contain the value of zero, indicates that the estimate of the indirect effect is not due to random chance. From a theoretical perspective, this finding provides support for the integration of AI governance-related risks into the VAM. VAM proposes that users evaluate technologies by comparing perceived benefits with perceived sacrifices. In the context of the present study, EUA represents an important perceived sacrifice because users may feel that AI applications reduce their ability to control decisions or influence outcomes. However, the mediation result demonstrates that this sacrifice can have broader consequences. When users experience reduced agency, they may become more sensitive to whether the AI system is transparent and understandable. If the system is perceived as opaque, users may have difficulty evaluating the fairness of its decisions, which can subsequently increase perceptions of inequity and discomfort. It follows that since both the lower and upper limits of the confidence interval are above zero, and the value of zero does not lie within the range of the confidence interval, then it means that there is indeed a significant difference between the indirect effect and zero. It further offers additional proof regarding the significance of the mediation effect. In mediation analysis, where there is no value of zero in the confidence interval, it implies that the indirect effect estimate is not just coincidental. Theoretically, this implies that there is indeed justification for integrating AI governance-related risks into the VAM. VAM suggests that users adopt technology on the basis of comparison between perceived benefit and sacrifice. In relation to this research, the sacrifice of EUA is a significant one since users might believe that the use of AI has made them less capable of controlling the situation and influencing the outcome. Nevertheless, as the mediation analysis shows, this sacrifice could have some wider implications. As users' agency is lowered, they might become more susceptible to the transparency and intelligibility of the system itself. As a result, if the system is not transparent for them, they will find it hard to assess the fairness of the decision-making process. The result could also be explained using AI transparency and accountability as the framework. User agency and AI transparency go hand in hand. The sense of agency is felt by the user when he or she knows the reasons behind the decisions of the AI and has the opportunity to question or alter the decisions. On the other hand, an AI which acts as a "black box" does not allow the user to know the reasons behind any decisions, recommendations or classifications made by the AI. This conclusion is further significant from the standpoint of fairness perception. User evaluation of algorithmic fairness does not necessarily have to be limited only to the outcome. Their perception may also be conditioned by their understanding of the procedure by which this outcome is reached. In the case of transparent and explainable AI decisions, users will be able to determine the reasonableness of this decision better. On the other hand, in the case of an opaque procedure, users are likely to perceive unexplained decisions as unfavorable. Hence, the mediation pathway in Table 8 can help reduce the negative implications of low user agency.

There is another implication that deals with users' discomfort. Users' discomfort in relation to AI might arise when people feel that there is no proper explanation or human involvement in making a decision. In cases when an AI algorithm suggests something, or makes a decision about eligibility, assessment of performance, or some other type of decision,

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

the users might feel discomfort because of the absence of a proper explanation of the decision-making process. This might be especially true if users do not feel they have enough power to affect the AI-assisted process at all. The results from mediation analysis are also applicable in the area of AI governance in a responsible manner. An organization would not be able to address the concerns of the users by providing them only with some level of control. If an organization offers the controls to users but lacks transparency, their problems might still persist. Thus, a good governance strategy must consist of control, transparency, explainability, accountability, and fairness elements. It is desirable that the users know how the system works, why certain decisions have been made by the AI, what information was used for that, and what actions the user can take. This is especially applicable to the topic under investigation since it deals with the usage of AI in devices and applications employed for personal and professional purposes. On a personal level, users will come across AI in recommendation systems, virtual assistants, customized products, smart devices, and other consumer products. At a professional level, AI can be used in the process of hiring, performance appraisal, purchase, decision making, and many others. In both instances, users may feel that their agency is compromised because of decisions or recommendations made by the AI system which cannot be easily questioned or understood. The conclusion indicates that the issue of opacity of algorithms does not represent only the problem associated with technology itself. On the contrary, it can be interrelated with the issue of users' perception of control and autonomy and their evaluation of fairness and satisfaction from the use of AI systems. The matter is especially significant in that users might have more tolerance to the level of automation if they feel that they have control over the process and perceive it properly. This finding even further solidifies the rationale for implementing a governance structure for AI that puts consumers at the center of the process. A governance structure that focuses on consumers realizes that a responsible AI process does not only revolve around the technical capabilities of the AI technology or cybersecurity or regulations.

It takes into account the experience of the users in using AI technologies. The strong indirect effect indicates that the experience of the users is interrelated. The results from testing for the mediation effect would also bring an extra benefit to the empirical contribution made by the research. From the results found in the direct effects of EUA, AOR, and PAID, one finds that all three are significant individually. However, the results of the mediation effects provide more information about the relationship among these variables. This is because the results provide an adequate explanation of the path from Erosion of User Agency → Algorithmic Opacity Risk → Perceived Algorithmic Inequity and Discomfort. To put into practice, organizations that build and apply AI technologies should offer clear algorithmic decision explanations, understandable privacy and use of information about data, user control, and review opportunities. Thus, people may feel less out-of-control in dealing with AI. Furthermore, transparency is likely to increase the ease of evaluating algorithmic results, decreasing feelings of unfairness and discomfort. The mediation result, however, needs to be viewed as an indication of an indirect statistical relationship and not conclusive evidence of causality, especially when the research makes use of cross-sectional survey data. The mediation result shows that the collected data are consistent with the theoretical framework of the mediation process, but further longitudinal and experimental studies would have helped determine the temporal relationship between the constructs of user agency, algorithmic opacity, and perceptions of inequity and discomfort. A number of arguments are provided from the data shown in Table 8 that AOR serves as the mediator between EUA and PAID. The positive indirect effect as well as the significance level of $p < 0.001$ show that reduced user agency causes increased perception of algorithmic opacity, and this in turn contributes to inequity and discomfort. The confidence interval (0.1 1, 0.1 2) reported for the indirect effect at the 95% level of significance confirms the significance of the indirect effect since zero is not included in the interval. Thus, the mediation hypothesis is confirmed. This result emphasizes the need for transparency, explainability, user control, and accountability in AI governance.

It can be concluded that handling algorithmic opacity can become an effective way to address negative user perceptions of AI products. What this means in simple terms is that the lower the levels of control that the users perceive, the larger the perceived risks of opacity of the algorithm, and subsequently (Laux, 2024), the larger the perceived inequity and unease in the algorithmic interaction (Subías-Beltrán et al., 2025). Although this indirect effect is found to be statistically significant based on the provided results, supported by the high significance of the findings ($p < 0.001$), it needs to be noted that the 95% confidence interval ranges from -0.482 to 0.504, thereby making it uncertain what exactly the findings have been able to ascertain in this respect (Hong et al., 2014). This view chimes with socio-technical thinking:

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

transparency is not separable from users' ability to participate and to challenge what they are shown (Recker et al., 2019). When individuals perceive CCER, opacity tends to feel more real; cybersecurity concerns amplify the sense that how algorithms decide things and how those decisions are secured is unclear (Beerends and Aydin, 2025). However, the direct influence of IIR and CCER on the perceived unfairness or justness of algorithmic output (Zanotti, 2025), or on how uncomfortable they feel about them, is weak or not statistically significant (Krokowski and Hirsch-Kreinsen, 2025). In a nutshell, people do not instinctively interpret privacy or cybersecurity risk as unfairness in itself (Thomas et al., 2025). Rather, these risks appear to inform broad conceptions about algorithmic transparency that, in turn, have more indirect effects on fairness perceptions (von Zahn et al., 2025). The corrosive effects of the lack of user agency are revealed to be the key risk factor that "nudges" users toward perceptions of unfairness (Shiau et al., 2024), highlighting the importance of perceptions of control/constraint, or having "a voice," in our determination of fairness (Choi and Johnson, 2021). It is interesting that the risk of opacity in the algorithms is revealed to be a key, important risk factor that drives perceptions of inequality caused by the algorithms (Gwebu et al., 2020). It would seem that, if nothing else, perceptions of the lack of transparency feeding into the algorithms cause users to be worried about bias (Ye et al., 2024), inequality, or the unforeseen consequences thereof, even if the justification of such occurrences is not immediately apparent (AlGhamdi et al., 2020). Mediation analysis suggests that this threat of opacity was one "component" through which the erosion of user agency translates into perceptions of inequality (Fernando and Dawson, 2009), though the margin we are dealing with prevents broad over-analysis of this component (Puthal et al., 2025). Overall, the data suggest that, from the perception/behavioral side (Rau, 2004), inequality related to the algorithms causes more problems related to how the system is managed than to the direct reaction to risk (Kooper et al., 2011). It is the filter through which perceptions related to privacy, security, or control are necessarily refracted. The paragraph above outlines the way in which the current research can be applied in expanding the understanding of the VAM. This paper does not only highlight the fact that consumers' perception of AI risks decreases their adoption of the technology. Instead, it demonstrates that there is a more complex psychological process through which consumers assess AI.

It involves the assessment of the level of control consumers feel over the AI system, the transparency of the AI technology, and fairness and trustworthiness. The results expand the scope of the VAM to the area of AI governance through the inclusion of governance-related risks in the calculation of value. Based on VAM, consumers assess whether to adopt a technology based on their assessment of the relative weight of its perceived benefits (e.g., usefulness, convenience, efficiency) and sacrifices or costs (e.g., monetary cost, effort required, privacy risk). Nonetheless, the current literature on VAM has paid little attention to the governance-related issues that are distinctive to artificial intelligence, such as transparency, explainability, accountability, fairness, cybersecurity, and consumer autonomy. The present research contributes to the theory in that it shows that consumers weigh these governance-related risks against each other when making a decision about whether AI delivers enough value to be adopted. One especially important implication of the research is that users' sense of control acts as an antecedent; it is one of the first psychological variables affecting the consumers' perception of AI systems. A sense of control means that users feel themselves capable of affecting, monitoring, understanding or overriding the decisions made by the AI system. Users feel a lack of control when they believe that AI systems make decisions on their own, not allowing any kind of human interaction with the process of decision-making. Contrary to what we may expect from the lack of control over technology, according to the authors, it does not lead directly to the perception of AI as an unfair technology. First, it leads to a new perception of how transparent the technology is. The results from statistics confirm this connection through the observation that AOR has a higher value in the case of low levels of perceived control ($\beta = .228$). The standardized path coefficient ($\beta = .228$) shows that there is a positive correlation between the reduction of consumer control and higher levels of perceived opacity. In effect, lack of ability on the part of consumers to have an impact on and comprehend the functioning of the AI algorithm creates a feeling among them that the system is a "black box," whose mechanisms are unknown to the users. Moreover, the findings of the study show that algorithmic opacity operates as a mediator, that is, as a psychological variable that explains the process by which one thing leads to another. The mediator variable describes the process through which reduced consumer control impacts their perceptions of AI being opaque and biased. Unlike a simple process, in which reduction in control directly impacts perceptions of bias, the results reveal a complex process, namely that reduced

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

control enhances consumers' perceptions of AI being opaque, and that the latter impacts perceptions of bias and undesirable decisions made by the algorithm. This is theoretically significant since it indicates that consumers do not directly construe governance risks, such as privacy risks and cybersecurity risks, as being indicative of unfairness.

Instead, the governance risks must first be construed under the prism of transparency. For instance, in the case where consumers perceive AI systems as collecting too much private data or making automatic decisions without giving the user any option to make input, this does not mean that these AI systems are unfair. Instead, what happens is that these consumers will first consider whether they understand how the algorithm functions and why the algorithm made certain decisions. When the algorithm seems to lack transparency in explaining how certain decisions were made, then they will find the decisions to be unfair. According to the research, the indirect effect is also statistically significant ($p < 0.001$). Being statistically significant, this term means that the mediation effect identified in the study is highly unlikely to occur due to mere chance. A p-value less than 0.001 means that the probability of observing such an indirect relationship is less than 0.1%, which proves that the mediation effect is indeed caused by algorithmic opacity and not random fluctuations. In other words, the data confirm the proposed theoretical model in the research. On the other hand, the confidence interval gives reason to interpret the results with caution. The confidence interval can be understood as a range within which the true effect for the population will fall. When the confidence interval is quite broad or contains values that diminish accuracy, then the true effect might not be as precise as required. In other words, even though the mediation effect proves to be statistically significant, the authors of the research recommend interpreting the size of the indirect effect cautiously and conducting future studies using larger samples. In general, the results of the research provide theoretical and practical contributions through a better understanding of the way consumer's mind work in relation to AI governance risks. Unlike other studies that suggest that privacy risk, cybersecurity risk, or loss of autonomy directly affects trust and acceptance of AI technologies, the authors have shown that these types of risks first impact the perception of transparency, which then impacts perceptions of fairness and emotions. The chain of "loss of control → increase in algorithmic opacity → perception of inequity" is the better way to understand how AI is accepted from the psychological point of view. Consumers judge AI technologies through a series of judgments. Therefore, the present research contributes to AI governance theory by uncovering the important psychological process through which governance-related risks affect consumer trust and acceptance. The results of this study not only contribute to the theory but also provide useful suggestions for the designers and implementers of AI technologies, who can be assured that increased transparency, explainability, and user control can result in decreased opacity and improved trust in AI technologies.

Theoretical implications

This current study contributes both theoretically and empirically to the extant literature in several significant ways on issues of AI, information governance, algorithmic transparency, and technology adoption. Through the incorporation of risk factors related to governance in the VAM, this study provides additional insight into how the consumer assesses AI systems. The implications of this study for the development of theories on technology adoption, AI information governance, and the psychological mechanisms that underlie consumers' evaluation of transparency and governance of AI-based systems are discussed. Firstly, the current research broadens the VAM by incorporating the concept of AI governance risks, which are significant value sacrifices that affect consumer evaluation of AI technologies. As per the traditional approach of the VAM, technology adoption is explained by suggesting that people make a comparison between the benefits and sacrifices associated with technology before making adoption decisions. Previous instances of using VAM have generally considered these sacrifices as economic cost, effort expectation, technological complexity, performance uncertainty, and privacy issues. However, while these concepts are applicable in some ways, they do not completely represent the unique risks of governance involved with modern AI technologies. Contrary to normal information technologies, AI technologies involve autonomous decision-making, continuous collection and processing of personal information, learning ability of machines, and functioning on complex algorithms that cannot be understood easily by regular users. Hence, when evaluating AI technologies, consumers take into consideration not only the utility benefits but also governance-related issues such as collection of their personal information, security of AI technologies, control over automation, and transparency of algorithms.

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

In order to fill this gap in theory, this paper will integrate IIR, CCER, EUA, and AOR as governance-oriented sacrifices in the VAM framework. The extension of the VAM shows that the value judgments by consumers in the context of AI are determined not only by economics but also by governance-related issues, which are not limited to economics and technology at all. The results of this study allow broadening the explanatory power of VAM, illustrating that responsible governance has become an integral part of value in the AI context. Second, the research has made a major contribution towards the theory of information governance in the area of AI by changing the perspective of the information governance framework from organizational compliance to consumer engagement and perceptions about information governance. Information governance research has so far been mainly centered around the responsibilities and roles of institutions like regulations, law, data governance policies, information security management, and organizational controls. While these perspectives are essential for responsible use of AI technology, they have often analyzed the effectiveness of the information governance practices through the lens of organizations and regulators and not through the lens of the consumers. On the contrary, the concept of governance effectiveness needs to be assessed in terms of the perception of consumers regarding the use of AI technologies. It is obvious that the measures will achieve their intended purpose when consumers believe that AI technologies are being applied in an open, safe, autonomous, and fair way. The consumer-centric perspective provides the possibility to extend the scope of the information governance model beyond organizational context due to the recognition of the fact that the effectiveness of governance relies on users' understanding, engagement, trust, and acceptance. The research has made a major contribution toward the theory of information governance in the area of AI by changing the perspective of the information governance framework from organizational compliance to consumer engagement and perceptions about information governance. Information governance research has so far been mainly centered around the responsibilities and roles of institutions like regulations, law, data governance policies, information security management, and organizational controls. While these perspectives are essential for responsible use of AI technology, they have often analyzed the effectiveness of the information governance practices through the lens of organizations and regulators and not through the lens of the consumers. On the contrary, the concept of governance effectiveness needs to be assessed in terms of the perception of consumers regarding the use of AI technologies.

It is obvious that the measures will achieve their intended purpose when consumers believe that AI technologies are being applied in an open, safe, autonomous, and fair way. The consumer-centric perspective provides the possibility to extend the scope of the information governance model beyond organizational context due to the recognition of the fact that the effectiveness of governance relies on users' understanding, engagement, trust, and acceptance. Third, the findings make an important theoretical contribution to the literature on algorithmic transparency and algorithmic opacity by identifying algorithmic opacity as a central governance construct that links AI-related risks to fairness perceptions. Previous studies have consistently argued that transparency is essential for promoting trust, accountability, and fairness in AI systems. However, much of this literature has remained conceptual or normative, focusing primarily on the importance of explainable AI without empirically demonstrating the mechanisms through which transparency influences consumer perceptions. The current study contributes to the literature on this topic in that it empirically establishes algorithmic opacity as a critical psychological construct underpinning the evaluation of AI governance. In the structural model proposed in the study, IIR, CCER, and EUA have been shown to strongly influence consumers' perception of algorithmic opacity, which in turn affects their perception of algorithmic inequity and discomfort. From the results of this study, it can be concluded that algorithmic opacity is more than a technical feature of AI; instead, it is a governance construct through which consumers interpret legitimacy, accountability, and fairness of the automated decision process. Furthermore, this research contributes to the academic literature on algorithmic fairness by disputing the conventional theory according to which the risks related to privacy and cybersecurity issues directly define consumers' fairness perceptions. Many studies tend to state that people who perceive high risks associated with privacy or cybersecurity are bound to see AI systems as unfair. The results of this research, on the contrary, provide evidence of a complex cognitive process. IIR and CCER have a significant impact on the perceptions of algorithmic opacity; at the same time, their influence on perceptions of algorithmic inequity and discomfort is mainly mediated. That means that people do not automatically associate their concerns regarding data collection and cybersecurity with unfair decisions made by algorithms; first, they assess whether AI systems are clear and comprehensible for consumers and only then they make

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

their fairness judgment. Thus, the research makes an important theoretical contribution since it adds a new layer to the understanding of the relationship between governance risks and algorithmic fairness perceptions; it becomes evident that consumers' fairness perceptions are mediated not only by risks but also by the quality of governance.

Fourth, in terms of contributions to governance theory, the study highlights the mediating influence of algorithmic opacity on the effect that the erosion of agency has on consumers' perceptions of unfairness and discomfort caused by algorithms. The results of the mediation test suggest that consumers who feel that their agency over AI is declining first experience increasing feelings of opacity, which later translates into increased perceptions of unfairness and discomfort. The finding adds support to governance theory that focuses on explainability, accountability, and human control as key elements of good governance of AI. At the same time, it also shows that consumer perceptions of the fairness of the technology are shaped by interrelated psychological mechanisms and not by single direct effects. Finally, through its interdisciplinary approach and empirical model linking technological risks, governance perceptions, and consumer-related impacts, this study contributes to AI governance, information systems, consumer behavior, and technology adoption literature. Prior research has tended to consider privacy, cybersecurity, transparency, fairness, and user control as separate themes, leading to disconnected theoretical perspectives on the phenomena of interest. Instead, this paper considers how these constructs can be combined into one theoretical framework explaining how a variety of risks connected with governance affect consumers' perceptions of AI-based products. This approach is important for developing new theory and showing that, in fact, AI governance must be considered as a multifaceted phenomenon, involving many technological, psychological, organizational, and governance-related aspects influencing consumer perceptions. Overall, this study adds to the wider theoretical body of literature on human-centric AI insofar as the effectiveness of AI governance is predicated on both good technology and the ability to foster positive consumer perceptions regarding the transparency, accountability, and empowerment of these AI systems. As per the results of the current study, the consumers do not judge AI systems based on their technical sophistication, but according to governance-related principles, such as transparency, controllability, etc. Technically advanced AI systems that lack transparency and controllability on the part of consumers tend to be perceived negatively by them. On the contrary, systems that allow for transparent explanation, user control, and transparent governance tend to be viewed as fair and trustworthy. Hence, there is a rising theoretical notion that the responsibility of an AI system should be considered in terms of humanistic values and principles. In conclusion, this paper contributes to theoretical knowledge in many ways, by expanding the scope of the VAM to cover AI governance risks, contributing to information governance theory from a consumer perspective, empirically confirming the importance of algorithmic opacity, developing the theories of algorithmic fairness from the viewpoint of a layered cognitive process, and integrating all risks related to governance into one model that explains consumers' perception of AI. Theoretical contributions presented here form the basis for future research on responsible AI adoption.

Practical implications

The results of this study point to a few concrete takeaways for those building AI, the platforms hosting it, and the policymakers shaping information governance. First, there is the clear implication that one of the linchpins to assuaging concerns about fairness is reducing the opacity of algorithmic reasoning. Investing in explainable AI, clear disclosures of data use, and decision rationales intelligible to humans goes a long way toward mitigating consumer anxiety, even if risks cannot be reduced to zero. Second, increasing user control is important. When users can challenge decisions, see exactly how their data are being used, and have a real say over AI-driven processes, feelings of inequity decrease and, subsequently, so do concerns about opacity. Therefore, user dashboards, appeal mechanisms, and hands-on, human-in-the-loop governance models help provide clarity. Third, privacy and cybersecurity plans must not be considered exclusive items that provide merely technical safety. Rather, they should be viewed as tools that provide transparency. When we clearly spell out who has access to users' data, data retention periods, etc., we can reduce the perception that the system is not transparent. Finally, regulators and policymakers must include perception metrics in AI governance. Because, even today, the traditional, outdated, time-tested compliance models will not understand the perception issue. Moreover, the "push factors" like explainability, accountability, and public participatory governance, can help increase trust and reduce perceptions of unfairness, etc. The authors come to the conclusion that explainability and transparency

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

are crucial governance mechanisms rather than merely technical characteristics. Also, organizations must move from strictly internal, compliance-based governance to consumer-centric tactics that enhance system explainability and reinstate a sense of user agency in order to promote adoption and build trust. Regarding how this research fits within the existing literature, the authors point out that the organizational and regulatory aspects of AI are the main topics of the literature currently being published. By focusing on the consumer as an active, logical decision-maker within the framework of the VAM, this study sets itself apart. Additionally, while many studies discuss the “AI black box” as a technical challenge, this research repositions algorithmic opacity as a central governance mechanism. It argues that opacity is the primary link through which governance failures are translated into negative consumer perceptions of fairness and justice. In this manner, this research helps software developers by providing recommendations and guidance. Finally, the study offers a value-based, as opposed to strictly control-based, viewpoint by incorporating VAM, a model commonly used for technology adoption, into the context of AI governance. It confirms that individual users' perceptions of value relative to the loss of control, privacy, and security are just as important to consumer trust as institutional protocols. The authors state that there are certain aspects of the manuscript presentation and research process that could have been improved, especially in relation to the practical perspective on the issues and the significance of the findings for governance. This is important as refining the presentation and development of the arguments will improve the manuscript's contribution to the scholarly discussion and overall understanding. It should be acknowledged that the authors agree with the idea that the work brings value to the current literature on the topic of AI governance.

Limitations

Nevertheless, it is essential to point out that this paper is not without some limitations, which need to be noted. For instance, the results obtained may be limited by the consideration that the survey is cross-sectional, which limits the assertion of the causal relationships between the variables under consideration. While the structural model produced a reasonable fit with the data, it is important to point out here that the findings might have better reflected the data if causality had been established, especially if the survey had been longitudinal or experimental in nature. Moreover, the results collected would best be described as the sentiments of people largely comprised of the younger generation and, more importantly, the more educated, considering the preponderant number of students and other learned people surveyed, which may not be the case among the general population and those with lower levels of education, especially considering the risks as seen by such people. Third, this research relies on people's own perceptions rather than hard objective measures of how AI systems actually behave, how transparent they are, or how fair they are. Thus, the findings reflect subjective views of governance risks rather than verified technical properties of AI applications. This perceptual focus, while theoretically defensible, could invite common method bias or social desirability effects. Fourth, despite the fact that the mediation analysis indicates a meaningful indirect role for algorithmic opacity, the confidence interval for that indirect effect includes zero - a signal that there is uncertainty as to how robust the mediation really is. Alternatively expressed, the proposed mediating mechanism, though theoretically embedded, requires additional empirical support. Finally, the research investigates only a thin slice of governance-related risks. Other important factors, such as trust in institutions, awareness of regulations, cultural norms, or perceived AI benefits, were not included and may interact with the studied constructs to form consumer judgments.

Future research framework

Based on the empirical findings of this study, future work must extend the theoretical boundaries of the proposed framework to create a more comprehensive and multidimensional framework for the governance of AI. While the current study attempts to shed light on how consumers view some critical risks of AI, namely, IIR, CCER, EUA, AOR, and PAID, future work can build on this framework and embed these factors within a broader ecosystem consisting of technology, organization, regulation, and society. Future research can thus expand the range of antecedents investigated in the current study by considering the challenges faced by consumers with respect to technology and governance arising from the fast-changing realm of AI. With AI developing at an increasing pace through the advancement of generative AI, autonomous AI, explainable AI (XAI), foundation models, multimodal AI, and AI agents that are capable of making decisions on their own, consumers will have to face novel forms of uncertainty and risk. Technological developments will give rise to ever-growing complexities concerning the issues of data privacy, algorithmic transparency, cybersecurity,

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

intellectual property rights, misinformation, bias, surveillance, and accountability. It is likely that informational intrusion and cyber exposure considered in this study as the dimensions of consumer risk will become more broadly defined governance risks. From an individual perspective, future studies can continue examining the psychological and behavioral factors that affect consumers' reactions to artificial intelligence systems. The psychological factors of trust, manipulation perception, autonomy, privacy protection, transparency, fairness, and control may play a role in determining how consumers perceive decision-making processes that are conducted by algorithms. Factors like AI literacy, digital competency, prior use of AI systems, perceived usefulness of AI, perceived intelligence of AI, emotional bond to AI systems, technology anxiety, and privacy consciousness may have significant impact on how consumers perceive algorithmic opacity and the fairness of decisions made by AI. By examining such psychological factors, it will become possible for researchers to explain why different AI systems are viewed in different ways by people with different knowledge and experience levels.

In addition to looking at the consumer perspective, the future research should look at some other variables related to organizations which may be acting as mediators between AI technology and consumer experience. Organizations play a key role in ensuring there is proper governance that impacts how the consumers perceive issues of transparency, accountability, and ethical use of AI. Variables like organization maturity for AI governance, ethical leadership, internal policies for AI, data governance approach, initiatives for transparency, explanation, human involvement, auditing, and responsible AI usage might act as mediators and amplify or reduce the effects of AI technology risk perception. In the same way, the institutional environment is yet another critical research field in which more exploration can be carried out. Factors such as regulatory maturity, effectiveness of enforcement, legal accountability, industry standards, cultural values, national AI policies, and governmental policy can affect the perception of AI governance by the consumers. This is because of the huge difference in the regulatory approach from country to country and even from industry to industry, making it possible to conduct cross-national researches and compare the effect that different governance structures have on the perception of AI. Other factors that also need to be tested in future research studies include moderation variables that would impact the strength and nature of the relationships discovered through this study. Factors include AI trust, institutional trust, digital literacy, AI experience, technological self-efficacy, privacy concern, risk-taking behavior, reputation of the organization, regulatory environment, cultural orientation, age, education, and industry context. This is because individuals who have higher awareness regarding AI and are digitally literate may not feel threatened as much because of the opacity of the algorithm since they have greater confidence in the technology. On the contrary, individuals who lack technological awareness and have higher privacy concerns may feel threatened and perceive more inequity. This is the case with firms that are operating in a highly regulated environment. One other fruitful avenue would be to look into additional outcome variables that may have been overlooked in this current research.

As it is, the research in question looks into the effect of consumer perception about algorithmic inequity and discomfort as outcomes. However, future research may look into more governance related outcomes, which may include consumer trust, acceptance of AI, adoption and continuance intentions, perception of value, legitimacy, organization's reputation, satisfaction and loyalty of customers, willingness to share personal information, compliance with AI recommendations, resistance to AI, and engagement with AI-enabled products/services. Further, performance indicators of AI governance such as perceived accountability, transparency effectiveness, ethical compliance, confidence in regulation and governance quality may be considered. These future developments also present opportunities to extend the theoretical foundation of the VAM. While the current study conceptualizes perceived AI risks as value sacrifices that influence consumer evaluations, future research could integrate governance-related constructs as additional sources of both perceived benefits and perceived sacrifices. Aspects of governance, such as transparency, explainability, accountability, fairness, privacy, and regulation, may help to increase the value perceived through increasing consumer confidence in AI applications, while governance failures may decrease the perceived value through increasing the uncertainty and perceived risk. With such a comprehensive approach, scholars will be able to study the impact of governance on not only perceptions of value and intentions to adopt AI technologies but also trust and acceptance. Methodologically, it is imperative that future research adopt different research designs in order to improve on causal inference and enhance the generalizability of results. With longitudinal research designs, researchers will be able to assess how consumer perceptions about AI governance change with increased exposure of the individual to the AI technologies and

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

development of the framework of governance. Researchers can use experimental research designs by manipulating variables such as transparency, explainability, disclosures of fairness, and human oversight in order to develop a causal relationship between governance and consumer perceptions. Moreover, any future studies on the subject should combine perceptual views of consumers with objective indicators of governance to construct more elaborate frameworks for AI governance.

The latter can involve algorithm explainability and auditability indexes, fairness measurements, conformity to international standards of AI governance, transparency reporting, cybersecurity measurements, and reviews carried out by regulatory bodies or independent auditors. A combination of perceptual and governance indicators would yield better empirical evidence on the efficiency of AI governance mechanisms, overcoming the shortcomings of survey-based self-reporting. Firstly, any future research on the topic would greatly benefit from taking the multi-stakeholder approach into account, thereby considering the perspectives of consumers, AI creators, organizational managers, regulators, policymakers, technology vendors, auditors, and NGOs. Indeed, AI governance is inherently multi-disciplinary in nature, with various technical, legal, ethical, management, and social aspects that can hardly be covered from the standpoint of consumers alone. The study of such perspectives would allow developing a much more comprehensive framework for AI governance that would help tackle the complexity of the issues related to responsible AI development and deployment. To illustrate this, there are various situations where individuals might have some experience working with AI, might be digitally literate, might be able to regulate trust, or have other factors where the impact of perceived inequality might differ. The structure of laws in various countries or industries might affect perceptions about governance. Adding moderators will improve not only how well such models can be explained in the future but also how sensitive these can be to a given situation. For instance, in terms of future studies, it will be possible to extend the outcome layers not only to perceived algorithmic unfairness and discomfort, as highlighted in previous works, but also to include different perspectives on what people perceive about system capabilities, such as trust factors, governance elements, and risk management indicators. In some respects, such a study can be synonymous with a value-based adoption model, as it will be possible to undertake a study on what governance has on not only different value judgments and ongoing use intentions but also enduring trust in AI applications. However, in order to adequately utilize such a framework as depicted in this study, it is critical to emphasize different longitudinal, experimental, and mixed-method designs that combine what individuals perceive with objective governance factors such as different "explainability" features within AI, as well as "auditability" metrics and factors from different regulatory bodies, with input from multiple stakeholders to provide a substantial basis for such a suggested framework.

Conclusions

This research investigated perspectives on information governance as they facilitate AI, guided by the VAM. The research sought an understanding of the role that risks concerning AI have in the non-transparency, inequity, or discomfort users feel about AI. The integration of IIR, CCER, loss of consumer control, and risk of algorithmic opacity into a single empirical model ensures an inherently consumer-centric approach to understanding AI information governance. The research revealed that IIR, CCER, and loss of consumer control all contributed to an increase in users' perceptions of non-transparency. The loss of control users feel about their interactions with AI decisions had a particularly significant influence, emphasizing that users being able to control, challenge, or even comprehend AI decisions is integral to AI transparency. However, in contrast to the potential approaches to cope with the challenges of informational intrusion and cyber exposure, there is not always a natural association drawn by consumers about how these related problems make them perceive that the system of algorithms is unfair or causes discomfort. What really determines fairness perceptions seems to be automatically linked to considerations of opacity in the algorithms. The relationship here is very strong and positive. There is a great association found in the study between the degree of opacity in the given system of algorithms and the perceptions of unfairness or discomfort. Opacity seems to be a very important means to govern perceptions of unfairness because "the mediation analysis supported that opacity indeed in part serves as a pathway by which lack of user agency influenced unfairness perceptions, emphasizing the importance of transparency in alleviating

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). *Cureus J Bus Econ* 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

user-based fairness-related governance concerns.” What we should really be looking at in terms of why people perceive inequality in terms of conversations around algorithms is not really about fairness or unfairness; rather, it is more aligned with how people perceive the system at a governance level.

Appendices

Questionnaire

Dear participant,

My name is Raheela Batool, and I am a Ph.D. research scholar at Pusan National University, South Korea. I am currently conducting my research for my doctoral dissertation, which focuses on “What factors do customers perceive about AI applications for personal and professional use?” My research aims to analyze the positive and negative factors that affect the perceptions of customers regarding AI applications for personal and professional purposes. I would like to thank you for taking this survey and contributing to the research. I assure you that all the research information will be kept strictly confidential in accordance with the norms of research ethics. Furthermore, this survey is completely anonymous, and there are no risks of disclosing any personal information. This survey will take approximately 8 to 11 minutes to complete. I would like to appreciate you for your time and effort.

Demographics

Gender

- Male
- Female

Age

- 10-20
- 21-30
- 31-40
- 41-50
- 50 above

Qualification

- Bachelors
- Masters
- Post-Graduation

Income (Monthly)

- 500USD-1000USD
- 1000USD-1500USD
- 1500USD-2000USD
- 4000USD-8000USD
- Others

Profession

- Student

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

- Researcher
- Engineer
- Private Employee/Businessman
- Artists
- Others

For each of the questions below, circle the response that best characterizes how you feel about the statement, where 5 = strongly agree, 4 = agree, 3 = neutral, 2 = disagree, and 1 = strongly disagree.

The self-administered questionnaire used for data collection is presented in Table 9 in the Appendix.

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

Measurement Items	Strongly Agree	Agree	Neutral	Disagree	Strongly Disagree
IIR					
1. If their personal information is being collected by applications based on AI in excessive amounts, I do not think their algorithm is clear to me.	5	4	3	2	1
2. If my personal information is used by the applications related to artificial intelligence in a way that I did not originally agree to, then I am having difficulty understanding the way by which the applications' decisions are made.	5	4	3	2	1
3. I am unsure who can access my personal data in AI applications, and this makes them seem non-transparent and opaque to me.	5	4	3	2	1
4. The longer AI applications store my personal data, the less transparent the data processing and decision-making procedures seem to me.	5	4	3	2	1

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

CCER					
5. Since artificial intelligence applications wield significant influence over my personal information, I am quite uncomfortable about the treatment of users by the algorithmic decisions.	5	4	3	2	1
6. Concerns regarding the exposure of my personal data because of hacking and misuse make me feel apprehensive regarding the fairness of decisions made by AI technology.	5	4	3	2	1
7. If the cybersecurity breakdowns happen in AI, I am concerned that their algorithmic result might harm me.	5	4	3	2	1
8. The risk that some unauthorized individuals may take advantage of AI systems gives me the concern that some users may be negatively affected by an inequitable AI-related decision.	5	4	3	2	1

How to cite this article:

EUA					
9. I strongly believe that because of the use of my own data in numerous AI applications, I have no control over the decisions that are made for me.	5	4	3	2	1
10. AI applications can make critical decisions without human interference. The decision process in these applications is opaque to me.	5	4	3	2	1
11. When personal data is extensively utilized, challenging decisions made by artificial intelligence algorithms become difficult.	5	4	3	2	1
12. The manner in which the AI applications utilize my personal data is not elucidated sufficiently so that I can comprehend their reasoning.	5	4	3	2	1
AOR					
13. In current AI applications, when my data is threatened by cyberattacks, I feel a lack of control over the decisions that could	5	4	3	2	1

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

discriminate against users.					
14. In current AI applications, when my data is threatened by cyberattacks it is hard for me to understand decisions that could discriminate against users.	5	4	3	2	1
15. Due to the weaknesses that exist in cyberspace and that are out of my control, I am uneasy with the way decisions are made by applications concerning me by artificial intelligence.	5	4	3	2	1
16. If AI systems are vulnerable to unauthorized access, I am concerned that the results of these algorithms might not always favor all individuals.	5	4	3	2	1
PAID					
17. In cases when AI is non-transparent, it may lead to unfairness due to hidden biases.	5	4	3	2	1
18. Because AI decision-making is opaque, AI-driven results could lead to an exacerbation	5	4	3	2	1

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

of existing social inequality.					
19. The black box nature of AI applications makes me suspicious of their potential unintended outcomes.	5	4	3	2	1
20. As a result of AI "black boxes," failures decrease my trust in the use of AI technologies.	5	4	3	2	1

TABLE 9: Self-administered questionnaire

Source: Authors

Additional Information

Author Contributions

All authors have reviewed the final version to be published and agreed to be accountable for all aspects of the work.

Concept and design: Raheela Batool, Muhammad Azlaan Zubair

Acquisition, analysis, or interpretation of data: Raheela Batool

Drafting of the manuscript: Raheela Batool

Critical review of the manuscript for important intellectual content: Raheela Batool, Muhammad Azlaan Zubair

Supervision: Raheela Batool

Disclosures

Human subjects: Consent was obtained from all participants in this study. Informed consent was obtained from all participants prior to participation. The study employed an anonymous survey design, and no personally identifiable or sensitive personal data were collected. Participation was voluntary, and the research posed minimal risk to participants. According to the policies of the authors' institution, this study was exempt from formal ethics committee/IRB review.

Animal subjects: All authors have confirmed that this study did not involve animal subjects or tissue.

Conflicts of interest: In compliance with [the ICMJE uniform disclosure form](#), all authors declare the following:

- **Payment/services info:** All authors have declared that no financial support was received from any organization for the submitted work.
- **Financial relationships:** All authors have declared that they have no financial relationships at present or within the previous three years with any organizations that might have an interest in the submitted

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

work.

- **Other relationships:** All authors have declared that there are no other relationships or activities that could appear to have influenced the submitted work.

Data Availability Statements

The datasets (and/or code) supporting this study are available from the corresponding author upon reasonable request.

Acknowledgements

The authors would like to thank Dr Arifeen for his encouragement, guidance, and valuable feedback during the course of this work. His suggestions and support are gratefully acknowledged.

References

1. Abraham R, Schneider J, vom Brocke J: [A taxonomy of data governance decision domains in data marketplaces](#). Electronic Markets. 2023, 33:23. [10.1007/s12525-023-00631-w](#)
2. Abraham R, Schneider J, vom Brocke J: [Data governance: A conceptual framework, structured review, and research agenda](#). International Journal of Information Management. 2019, 49:424-438. [10.1016/j.ijinfomgt.2019.07.008](#)
3. AlGhamdi S, Win KT, Vlahu-Gjorgievska E: [Information security governance challenges and critical success factors: Systematic review](#). Computers & Security. 2020, 99:102030. [10.1016/j.cose.2020.102030](#)
4. Almulihi AH, Alassery F, Irshad Khan A, Shukla S, Kumar Gupta B, Kumar R: [Analyzing the implications of healthcare data breaches through computational technique](#). Intelligent Automation and Soft Computing. 2022, 32:1763-1779. [10.32604/iasc.2022.023460](#)
5. Alreemy Z, Chang V, Walters R, Wills G: [Critical success factors \(CSFs\) for information technology governance \(ITG\)](#). International Journal of Information Management. 2016, 36:907-916. [10.1016/j.ijinfomgt.2016.05.017](#)
6. Balkan D, Akyuz GA: [Artificial intelligence \(AI\) and machine learning \(ML\) in procurement and purchasing decision-support \(DS\): a taxonomic literature review and research opportunities](#). Artificial Intelligence Review. 2025, 58:341. [10.1007/s10462-025-11336-1](#)
7. Baskaran V, Davis K, Bali RK, Naguib RNG, Wickramasinghe N: [Managing information and knowledge within maternity services: Privacy and consent issues](#). Informatics for Health and Social Care. 2013, 38:196-210. [10.3109/17538157.2012.735732](#)
8. Beerends S, Aydin C: [Negotiating the authenticity of AI: how the discourse on AI rejects human indeterminacy](#). AI and Society. 2025, 40:263-276. [10.1007/s00146-024-01884-5](#)
9. Brown WC: [IT governance, architectural competency, and the Vasa](#). Information Management & Computer Security. 2006, 14:140-154. [10.1108/09685220610655889](#)
10. Bulgurcu B, Cavusoglu H, Benbasat I: [Information security policy compliance: an empirical study of rationality-based beliefs and information security awareness](#). MIS Quarterly. 2010, 34:523-548. [10.2307/25750690](#)
11. Burton-Jones A, Volkoff O: [How can we develop contextualized theories of effective use? A demonstration in the context of community-care electronic health records](#). Information Systems Research. 2017, 28:468-489. [10.1287/isre.2017.0702](#)
12. Cao G, Duan Y, Cadden T: [The link between information processing capability and competitive advantage mediated through decision-making effectiveness](#). International Journal of Information Management. 2019, 44:121-131. [10.1016/j.ijinfomgt.2018.10.003](#)
13. Chernyshev M, Zeadally S, Baig Z: [Healthcare data breaches: implications for digital forensic readiness](#). Journal of Medical Systems. 2019, 43:7. [10.1007/s10916-018-1123-2](#)
14. Chi M, Zhao J, George JF, Li Y, Zhai S: [The influence of inter-firm IT governance strategies on relational performance: The moderation effect of information technology ambidexterity](#). International Journal of Information Management. 2017, 37:43-53. [10.1016/j.ijinfomgt.2016.11.007](#)
15. Chima T, Mkwinda E, Kumwenda S: [The need for health information management professionals in Malawi health facilities](#). Health Information Management Journal. 2023, 53:6-13. [10.1177/18333583231180772](#)

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

16. Choi SJ, Johnson ME: [The relationship between cybersecurity ratings and the risk of hospital data breaches](#). Journal of the American Medical Informatics Association. 2021, 28:2085-2092. [10.1093/jamia/ocab142](#)
17. Churamani N, Kara O, Gunes H: [Domain-incremental continual learning for mitigating bias in facial expression and action unit recognition](#). IEEE Transactions on Affective Computing. 2023, 14:3191-3206. [10.1109/taffc.2022.3181033](#)
18. D'Arcy J, Hovav A, Galletta D: [User awareness of security countermeasures and its impact on information systems misuse: A deterrence approach](#). Information Systems Research. 2009, 20:79-98. [10.1287/isre.1070.0160](#)
19. De Haes S, Van Grembergen W: [An exploratory study into IT governance implementations and its impact on business/IT alignment](#). Information Systems Management. 2009, 26:123-137. [10.1080/10580530902794786](#)
20. Donaldson A, Walker P: [Information governance—a view from the NHS](#). International Journal of Medical Informatics. 2004, 73:281-284. [10.1016/j.ijmedinf.2003.11.009](#)
21. Dutta A, Roy R, Seetharaman P: [An assimilation maturity model for IT governance and auditing](#). Information & Management. 2022, 59:103569. [10.1016/j.im.2021.103569](#)
22. Fernández Llorca D, Hamon R, Junklewitz H, et al.: [Testing autonomous vehicles and AI: perspectives and challenges from cybersecurity, transparency, robustness and fairness](#). European Transport Research Review. 2025, 17:38. [10.1186/s12544-025-00732-x](#)
23. Fernando JI, Dawson LL: [The health information system security threat lifecycle: An informatics theory](#). International Journal of Medical Informatics. 2009, 78:815-826. [10.1016/j.ijmedinf.2009.08.006](#)
24. Fischer-Preßler D, Marx J, Bunker D, Stieglitz S, Fischbach K: [Social media information governance in multi-level organizations: How humanitarian organizations accrue social capital](#). Information & Management. 2023, 60:103838. [10.1016/j.im.2023.103838](#)
25. Gandía JL: [Determinants of internet-based corporate governance disclosure by Spanish listed companies](#). Online Information Review. 2008, 32:791-817. [10.1108/14684520810923944](#)
26. Gaozhao D: [Flagging fake news on social media: An experimental study of media consumers' identification of fake news](#). Government Information Quarterly. 2021, 38:101591. [10.1016/j.giq.2021.101591](#)
27. Gegenhuber T, Mair J, Lührsén R, Thäter L: [Orchestrating distributed data governance in open social innovation](#). Information and Organization. 2023, 33:100453. [10.1016/j.infoandorg.2023.100453](#)
28. Giannouchos TV, Ferdinand AO, Ilangovan G, Ragan E, Nowell WB, Kum HC, Schmit CD: [Identifying and prioritizing benefits and risks of using privacy-enhancing software through participatory design: a nominal group technique study with patients living with chronic conditions](#). Journal of the American Medical Informatics Association. 2021, 28:1746-1755. [10.1093/jamia/ocab073](#)
29. González-Sendino R, Serrano E, Bajo J: [Mitigating bias in artificial intelligence: Fair data generation via causal models for transparent and explainable decision-making](#). Future Generation Computer Systems. 2024, 155:384-401. [10.1016/j.future.2024.02.023](#)
30. Gwebu KL, Wang J, Hu MY: [Information security policy noncompliance: An integrative social influence model](#). Information Systems Journal. 2020, 30:220-269. [10.1111/isj.12257](#)
31. Hammerschmidt T, Hafner A, Stolz K, Passlack N, Posegga O, Gerholz KH: [A review of how different views on ethics shape perceptions of morality and responsibility within AI transformation](#). Information Systems Frontiers. 2025, [10.1007/s10796-025-10596-0](#)
32. Harvey MJ, Harvey MG: [Privacy and security issues for mobile health platforms](#). Journal of the Association for Information Science and Technology. 2014, 65:1305-1318. [10.1002/asi.23066](#)
33. Hassandoust F, Akhlaghpour S, Johnston AC: [Individuals' privacy concerns and adoption of contact tracing mobile applications in a pandemic: A situational privacy calculus perspective](#). Journal of the American Medical Informatics Association. 2021, 28:463-471. [10.1093/jamia/ocaa240](#)
34. Hassidim A, Korach T, Shreberk-Hassidim R, Thomaidou E, Uzefovsky F, Ayal S, Ariely D: [Prevalence of sharing access credentials in electronic medical records](#). Healthcare Informatics Research. 2017, 23:176-182. [10.4258/hir.2017.23.3.176](#)
35. Hong W, Chan FKY, Thong JYL, Chasalow LC, Dhillon G: [A framework and guidelines for context-specific theorizing in information systems research](#). Information Systems Research. 2014, 25:111-136. [10.1287/isre.2013.0501](#)
36. Islam MS, Farah N, Stafford TF: [Factors associated with security/cybersecurity audit by internal audit function](#). Managerial Auditing Journal. 2018, 33:377-409. [10.1108/maj-07-2017-1595](#)

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

37. Jalali MS, Landman A, Gordon WJ: [Telemedicine, privacy, and information security in the age of COVID-19](#). Journal of the American Medical Informatics Association. 2021, 28:671-672. [10.1093/jamia/ocaa310](#)
38. Johnston A, Di Gangi P, Howard J, Worrell JL: [It takes a village: Understanding the collective security efficacy of employee groups](#). Journal of the Association for Information Systems. 2019, 20:186-212. [10.17705/1jais.00533](#)
39. Joia LA, Correia JCP: [CIO competencies from the IT professional perspective: Insights from Brazil](#). Journal of Global Information Management. 2018, 26:74-103. [10.4018/jgim.2018040104](#)
40. Joshi S, Bhattacharya S, Pathak P, Natraj NA, Saini J, Goswami S: [Harnessing the potential of generative AI in digital marketing using the Behavioral Reasoning Theory approach](#). International Journal of Information Management Data Insights. 2025, 5:100317. [10.1016/j.jjimei.2024.100317](#)
41. Kim YJ, Lee JM, Koo C, Nam K: [The role of governance effectiveness in explaining IT outsourcing performance](#). International Journal of Information Management. 2013, 33:850-860. [10.1016/j.ijinfomgt.2013.07.003](#)
42. Kim HW, Chan HC, Gupta S: [Value-based adoption of mobile internet: An empirical investigation](#). Decision Support Systems. 2007, 43:111-126. [10.1016/j.dss.2005.05.009](#)
43. Kooper MN, Maes R, Lindgreen EEOR: [On the governance of information: Introducing a new concept of governance to support the management of information](#). International Journal of Information Management. 2011, 31:195-200. [10.1016/j.ijinfomgt.2010.05.009](#)
44. Kordzadeh N, Warren J, Seifi A: [Antecedents of privacy calculus components in virtual health communities](#). International Journal of Information Management. 2016, 36:724-734. [10.1016/j.ijinfomgt.2016.04.015](#)
45. Krokowski T, Hirsch-Kreinsen H: [Once promised, forever believed? An essay on the blind spots of AI promises](#). AI & Society. 2025, 41:3497-3510. [10.1007/s00146-025-02765-1](#)
46. Laux J: [Institutionalised distrust and human oversight of artificial intelligence: towards a democratic design of AI governance under the European Union AI Act](#). AI & Society. 2024, 39:2853-2866. [10.1007/s00146-023-01777-z](#)
47. Lämmermann L, Hofmann P, Urbach N: [Managing artificial intelligence applications in healthcare: Promoting information processing among stakeholders](#). International Journal of Information Management. 2024, 75:102728. [10.1016/j.ijinfomgt.2023.102728](#)
48. Lin C, Wittmer JLS, Luo XR: [Cultivating proactive information security behavior and individual creativity: The role of human relations culture and IT use governance](#). Information & Management. 2022, 59:103650. [10.1016/j.im.2022.103650](#)
49. Majrashi K: [Employees' perceptions of the fairness of AI-based performance prediction features](#). Cogent Business & Management. 2025, 12:2456111. [10.1080/23311975.2025.2456111](#)
50. Mansoor M: [Citizens' trust in government as a function of good governance and government agency's provision of quality information on social media during COVID-19](#). Government Information Quarterly. 2021, 38:101597. [10.1016/j.giq.2021.101597](#)
51. Mikalef P, Boura M, Lekakos G, Krogstie J: [The role of information governance in big data analytics driven innovation](#). Information & Management. 2020, 57:103361. [10.1016/j.im.2020.103361](#)
52. Mueller ML: [It's just distributed computing: Rethinking AI governance](#). Telecommunications Policy. 2025, 49:102917. [10.1016/j.telpol.2025.102917](#)
53. Nexhip A, Riley M, Robinson K: [Defining career success: A cross-sectional analysis of health information managers' perceptions](#). Health Information Management Journal. 2025, 54:34-42. [10.1177/18333583231184903](#)
54. Nicho M: [A process model for implementing information systems security governance](#). Information and Computer Security. 2018, 26:10-38. [10.1108/ics-07-2016-0061](#)
55. Novelli C, Casolari F, Rotolo A, Taddeo M, Floridi L: [Taking AI risks seriously: a new assessment model for the AI Act](#). AI & Society. 2024, 39:2493-2497. [10.1007/s00146-023-01723-z](#)
56. Odilla F: [Unfairness in AI anti-corruption tools: main drivers and consequences](#). Minds and Machines. 2024, 34:28. [10.1007/s11023-024-09688-8](#)
57. Ogbanufe O, Kim DJ, Jones MC: [Informing cybersecurity strategic commitment through top management perceptions: The role of institutional pressures](#). Information & Management. 2021, 58:103507. [10.1016/j.im.2021.103507](#)

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

58. Olawade DB, Weerasinghe K, Teke J, Msiska M, Boussios S, Hatzidimitriadou E: [Evaluating AI adoption in healthcare: Insights from the information governance professionals in the United Kingdom](#). International Journal of Medical Informatics. 2025, 199:105909. [10.1016/j.ijmedinf.2025.105909](#)
59. Papagiannidis E, Enholm IM, Dremel C, Mikalef P, Krogstie J: [Toward AI governance: identifying best practices and potential barriers and outcomes](#). Information Systems Frontiers. 2023, 25:123-141. [10.1007/s10796-022-10251-y](#)
60. Passlack N, Hammerschmidt T, Posegga O: [With great power comes great responsibility: what shapes AI literacy for responsible interactions of knowledge workers with AI?](#). Information Systems Frontiers. 2025, 28:11-46. [10.1007/s10796-025-10648-5](#)
61. Paulus D, Fathi R, Fiedrich F, Van de Walle BV, Comes T: [On the interplay of data and cognitive bias in crisis information management: an exploratory study on epidemic response](#). Information Systems Frontiers. 2024, 26:391-415. [10.1007/s10796-022-10241-0](#)
62. Peng C, van Doorn J, Eggers F, Wieringa JE: [The effect of required warmth on consumer acceptance of artificial intelligence in service: The moderating role of AI-human collaboration](#). International Journal of Information Management. 2022, 66:102533. [10.1016/j.ijinfomgt.2022.102533](#)
63. Pool J, Akhlaghpour S, Fatehi F, Burton-Jones A: [A systematic analysis of failures in protecting personal health data: A scoping review](#). International Journal of Information Management. 2024, 74:102719-102719. [10.1016/j.ijinfomgt.2023.102719](#)
64. Potluri S: [Policy-aware secure data governance in distributed information systems using explainable AI models](#). International Journal of AI, BigData, Computational and Management Studies. 2025, 6:1-10. [10.63282/3050-9416.ijaibdcms-v6i3p101](#)
65. Puthal D, Mishra AK, Mohanty SP, Longo A, Yeun CY: [Shadow AI: cyber security implications, opportunities and challenges in the unseen frontier](#). SN Computer Science. 2025, 6:405. [10.1007/s42979-025-03962-x](#)
66. Rasouli MR, Trienekens JJM, Kusters RJ, Grefen PWPJ: [Information governance requirements in dynamic business networking](#). Industrial Management & Data Systems. 2016, 116:1356-1379. [10.1108/imds-06-2015-0260](#)
67. Rau KG: [Effective governance of it: Design objectives, roles, and relationships](#). Information Systems Management. 2004, 21:35-42. [10.1201/1078/44705.21.4.20040901/84185.4](#)
68. Recker J, Indulska M, Green P, Burton-Jones A, Weber R: [Information systems as representations: A review of the theory and evidence](#). Journal of the Association for Information Systems. 2019, 20:735-786. [10.17705/1jais.00550](#)
69. Ryan M, de Roo N, Wang H, Blok V, Atik C: [AI through the looking glass: an empirical study of structural social and ethical challenges in AI](#). AI & Society. 2025, 40:3891-3907. [10.1007/s00146-024-02146-0](#)
70. Samuel J, Kashyap R, Samuel Y, Pelaez A: [Adaptive cognitive fit: Artificial intelligence augmented management of information facets and representations](#). International Journal of Information Management. 2022, 65:102505. [10.1016/j.ijinfomgt.2022.102505](#)
71. Schneider J: [Explainable Generative AI \(GenXAI\): a survey, conceptualization, and research agenda](#). Artificial Intelligence Review. 2024, 57:289. [10.1007/s10462-024-10916-x](#)
72. Shiao WL, Chau PYK, Thatcher JB, Teng CI, Dwivedi YK: [Have we controlled properly? Problems with and recommendations for the use of control variables in information systems research](#). International Journal of Information Management. 2024, 74:102702. [10.1016/j.ijinfomgt.2023.102702](#)
73. Soomro ZA, Shah MH, Ahmed J: [Information security management needs more holistic approach: A literature review](#). International Journal of Information Management. 2016, 36:215-225. [10.1016/j.ijinfomgt.2015.11.009](#)
74. Subías-Beltrán P, Unceta I, de Lecuona I, Pujol O: [Asking the right questions: a governance approach to uphold human autonomy in artificial intelligence](#). AI & Society. 2025, 41:1149-1173. [10.1007/s00146-025-02611-4](#)
75. Tallon PP, Ramirez RV, Short JE: [The information artifact in IT governance: Toward a theory of information governance](#). Journal of Management Information Systems. 2013, 30:141-178. [10.2753/mis0742-1222300306](#)
76. Thomas C, Roberts H, Mökander J, Tsamados A, Taddeo M, Floridi L: [The case for a broader approach to AI assurance: addressing "hidden" harms in the development of artificial intelligence](#). AI & Society. 2025, 40:1469-1484. [10.1007/s00146-024-01950-y](#)
77. Tiwana A, Konsynski B, Venkatraman N: [Special issue: Information technology and organizational governance: The IT governance cube](#). Journal of Management Information Systems. 2013, 30:7-12. [10.2753/mis0742-1222300301](#)
78. Tiwari A, Farag HEZ: [Responsible AI framework for autonomous vehicles: addressing bias and fairness risks](#). IEEE Access. 2025, 13:58800-58822. [10.1109/access.2025.3556781](#)

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>

79. Vial G: [Data governance and digital innovation: A translational account of practitioner issues for IS research](#). Information and Organization. 2023, 33:100450. [10.1016/j.infoandorg.2023.100450](#)
80. von Zahn M, Zacharias J, Lowin M, Chen J, Hinz O: [Navigating AI conformity: A design framework to assess fairness, explainability, and performance](#). Electronic Markets. 2025, 35:24. [10.1007/s12525-025-00770-2](#)
81. Weber R: [Taking the ontological and materialist turns: Agential realism, representation theory, and accounting information systems](#). International Journal of Accounting Information Systems. 2020, 39:100485. [10.1016/j.accinf.2020.100485](#)
82. Xue Y, Liang H, Boulton WR: [Information technology governance in information technology investment decision processes: the impact of investment characteristics, external environment, and internal context](#). MIS Quarterly. 2008, 32:67-96. [10.2307/25148829](#)
83. Ye X, Yan Y, Li J, Jiang B: [Privacy and personal data risk governance for generative artificial intelligence: A Chinese perspective](#). Telecommunications Policy. 2024, 48:102851. [10.1016/j.telpol.2024.102851](#)
84. Zanotti G: [AI systems should be trustworthy, not trusted](#). AI & Society. 2025, 41:3401-3412. [10.1007/s00146-025-02728-6](#)

How to cite this article:

Batool R, Zubair M (September 14, 2026) An Empirical Analysis of Consumer Risk Perceptions in Artificial Intelligence (AI) Applications: AI-Driven Information Governance (AIG). Cureus J Bus Econ 3 : es44404-026-00278-8. DOI <https://doi.org/10.7759/s44404-026-00278-8>