

Term paper in seminar
Risk Perception and Communication

Does AI Regulation Increase Public Trust?

Risk Perception under the EU AI Act

Chair of Marketing and Consumer Research
Instructor: Prof. Dr. Jutta Roosen

Presented by: Katupulle Gedara Madushan [Matriculation No: 03795853]

Table of Contents

List of Tables.....	IV
List of Figures	V
 Chapter 1 : Introduction	 1
1.1 Introduction to the Study	1
1.2 Research Gap	2
1.3 Research Objectives.....	3
1.4 Research Questions	3
1.4.1 Main Research Question	3
1.4.2 Sub-Research Questions	3
1.5 Structure of the Paper.....	4
Chapter 2 : Literature Review	5
2.1 Artificial Intelligence and Public Trust	5
2.2 The European Union Artificial Intelligence Act	5
2.3 Perceived Risk in Artificial Intelligence	6
2.4 Adoption Intention	6
2.5 Theoretical Foundations.....	6
2.5.1 Risk Perception Theory.....	7
2.5.2 Signalling Theory.....	7
2.5.3 Trust Transfer Theory	7
2.5.4 Technology Acceptance Model	8
2.6 Previous Empirical Evidence	8
2.7 Hypotheses Development	9
Chapter 3 : Research Methodology.....	10
3.1 Research Design.....	10

3.2 Conceptual Framework.....	10
3.3 Experimental Stimuli and Procedure	10
3.4 Sample and Data Collection.....	11
3.5 Measurement and Operationalization of Variables	11
3.6 Data Analysis Methods	12
Chapter 4 : Results and Findings	13
4.1 Demographic Characteristics	13
4.2 Reliability Analysis and Descriptive Statistics	13
4.3 Correlation Analysis.....	14
4.4 Independent samples t-test of EU AI Act Disclosure and Trust in AI	14
4.5 Independent samples t-test of EU AI Act Disclosure and Perceived Risk.....	15
4.6 Regression analysis of Perceived Risk and Trust in AI	16
4.7 Regression analysis of Trust in AI and Adoption Intention	17
4.8 Mediation analysis of Perceived Risk.....	18
Chapter 5 : Discussion of Findings.....	21
5.1. EU AI Act Disclosure, Trust, and Perceived Risk	21
5.2 The Relationship Between Perceived Risk and Trust in AI	21
5.3 The Relationship Between Trust and Adoption Intention	21
5.4 The Mediating Role of Perceived Risk	22
5.5 Theoretical Implications	22
5.6 Practical Implications.....	23
Chapter 6 : Conclusion.....	25
6.1 Conclusion	25
6.2 Limitations and Future Research	25
Declaration on the Use of AI	27
References.....	28
Appendix.....	I

Appendix I : Survey Questionnaire.....	I
Appendix II : Control group (standard chatbot)	III
Appendix III : Treatment group (EU AI Act disclosure)	IV
Declaration of authorship.....	V

List of Tables

Table 3.1 Operationalization of Variables	12
Table 4.1 Demographic Characteristics of the Sample	13
Table 4.2 Reliability Analysis and Descriptive Statistics	14
Table 4.3 Correlation Matrix.....	14
Table 4.4 Independent Samples t-test for H1: EU AI Act Disclosure and Trust in AI.....	15
Table 4.5 Independent Samples t-test for H2: EU AI Act Disclosure and Perceived Risk.....	16
Table 4.6 Regression Results for H3: Perceived Risk and Trust in AI	17
Table 4.7 Regression Results for H4: Trust in AI and Adoption Intention	18
Table 4.8 Mediation Analysis Results for H5: The Mediating Role of Perceived Risk	19

List of Figures

Figure 3.1 Conceptual framework	10
Figure 4.1 Trust in AI by Experimental Condition	15
Figure 4.2 Perceived Risk by Experimental Condition	16
Figure 4.3 Effect of Perceived Risk on Trust in AI.....	17
Figure 4.4 Effect of Trust in AI on Adoption Intention	18
Figure 4.5 Mediation Model of the Relationship Between EU AI Act Disclosure, Perceived Risk, Trust in AI, and Adoption Intention.....	20
Figure 5.1 User Interpretation Mechanism of EU AI Act Disclosure as Protection and Risk Signals.....	23

Chapter 1 : Introduction

1.1 Introduction to the Study

Artificial Intelligence (AI) technologies are increasingly becoming integrated into various sectors of modern society, including healthcare, finance, education, transportation, manufacturing, and customer service. In recent years, advances in generative AI and conversational AI systems have significantly accelerated public interaction with AI technologies and expanded their applications in everyday life (Leschanowsky et al., 2024). As AI systems continue to influence decision-making processes and automate complex tasks, understanding the factors that shape public acceptance of these technologies has become increasingly important.

Despite the considerable benefits associated with AI adoption, concerns regarding privacy, fairness, transparency, accountability, and discrimination remain major barriers to public acceptance of AI systems (Leschanowsky et al., 2024). Similar concerns have accompanied previous technological revolutions throughout history. During the Industrial Revolution, workers feared displacement caused by mechanization, while later technological developments such as calculators, computers, and the internet generated scepticism regarding their societal consequences. Today, artificial intelligence has become the subject of similar debates concerning its impact on employment, privacy, ethical standards, and human autonomy.

Trust has emerged as one of the most important determinants of successful human–AI interaction and technology adoption (Scharowski et al., 2024). Individuals are more likely to rely on AI systems, accept AI recommendations, and integrate AI technologies into their daily activities when they perceive such systems as trustworthy and reliable. However, trust in AI is strongly influenced by perceived risks associated with data privacy, algorithmic bias, lack of explainability, and uncertainty regarding automated decision-making processes (Leschanowsky et al., 2024). Risk Perception Theory suggests that individuals evaluate risks subjectively rather than objectively and that these perceptions significantly influence attitudes and behavioral intentions toward technologies (Slovic, 1987).

In response to these concerns, policymakers around the world have increasingly emphasized the need for governance mechanisms that ensure the responsible development and deployment of AI systems. The European Union Artificial Intelligence Act (EU AI Act), adopted in 2024, represents the world's first comprehensive regulatory framework specifically designed to

govern artificial intelligence technologies (European Commission, 2024). The regulation adopts a risk-based approach and establishes requirements related to transparency, accountability, human oversight, and safety standards for AI applications based on their potential societal impact.

However, the effectiveness of regulatory disclosures in the context of AI remains uncertain. While policymakers intend regulations to reassure users, disclosures regarding AI regulation may also unintentionally increase awareness of potential risks associated with AI technologies. Generation Z represents a particularly relevant population for investigating these relationships. Individuals belonging to Generation Z have grown up in highly digital environments and are among the earliest adopters and most frequent users of AI technologies. As future employees, consumers, and decision-makers, their perceptions of AI are likely to influence the long-term societal acceptance and diffusion of artificial intelligence systems.

Therefore, this study investigates whether regulatory disclosure under the EU AI Act functions as a signal of safety that influences perceived risk and trust in AI systems among Generation Z users. Specifically, the study examines the relationships among EU AI Act disclosure, perceived risk, trust in AI, and adoption intention using an experimental research design.

1.2 Research Gap

Although previous studies have demonstrated that certifications, privacy labels, and regulatory mechanisms can positively influence trust in products and digital technologies, empirical evidence regarding the effectiveness of AI-specific regulations remains limited. Existing studies have primarily focused on privacy labels under the GDPR (Fox et al., 2022), certification labels for trustworthy AI (Scharowski et al., 2023), and the role of trust in human–AI interaction (Scharowski et al., 2024).

Furthermore, previous research has identified perceived risk as a major determinant of technology acceptance and trust formation (Slovic, 1987; Leschanowsky et al., 2024). However, little empirical research has examined whether disclosures related to the EU AI Act operate as signals of safety that reduce perceived risk and increase trust in AI systems. Similarly, the mechanisms through which regulatory disclosures influence trust remain largely unexplored.

In addition, Generation Z users have received relatively limited attention within the emerging literature on AI governance and regulation despite representing one of the most active user

groups of generative AI technologies. As a result, it remains unclear whether EU AI Act disclosures enhance trust by reducing perceived risks or whether they unintentionally increase users' awareness of potential dangers associated with AI systems.

Therefore, a significant research gap exists regarding the influence of EU AI Act disclosures on perceived risk and trust in AI systems and the mediating role of perceived risk within this relationship.

1.3 Research Objectives

The main objective of this study is to examine whether regulatory disclosure under the EU AI Act functions as a signal of safety that influences perceived risk and trust in AI systems among Generation Z users.

The specific objectives of the study are:

1. To examine whether EU AI Act disclosure influences perceived risk associated with AI systems.
2. To investigate the relationship between perceived risk and trust in AI systems.
3. To analyze whether EU AI Act disclosure influences trust in AI systems.
4. To examine the effect of trust in AI on adoption intention.
5. To determine whether perceived risk mediates the relationship between EU AI Act disclosure and trust in AI systems.

1.4 Research Questions

Based on the research objectives, this study addresses the following research questions:

1.4.1 Main Research Question

Can regulatory disclosure under the EU AI Act act as a signal of safety that reduces perceived risk and increases trust in AI systems?

1.4.2 Sub-Research Questions

1. Does EU AI Act disclosure influence perceived risk associated with AI systems?
2. How does perceived risk influence trust in AI systems?
3. Does EU AI Act disclosure influence trust in AI systems?
4. Does trust in AI influence adoption intention?

5. Does perceived risk mediate the relationship between EU AI Act disclosure and trust in AI systems?

1.5 Structure of the Paper

The remainder of this paper is organized as follows. Chapter 2 presents the literature review and theoretical foundations underlying the study, including Risk Perception Theory, Signaling Theory, Trust Transfer Theory, and the Technology Acceptance Model. Chapter 3 describes the research methodology, including the conceptual framework, experimental design, sampling strategy, measurement instruments, and analytical techniques employed in the study. Chapter 4 presents the empirical findings and hypothesis testing results. Chapter 5 discusses the findings in relation to the existing literature and highlights the theoretical and practical implications of the study. Finally, Chapter 6 concludes the study by summarizing the main findings, acknowledging limitations, and providing recommendations for future research.

Chapter 2 : Literature Review

2.1 Artificial Intelligence and Public Trust

Artificial Intelligence (AI) refers to computational systems capable of performing tasks that typically require human intelligence, including learning, reasoning, and decision-making (OECD, 2019). Recent advances in machine learning and generative AI have accelerated the adoption of AI technologies across various sectors, increasing the importance of understanding public attitudes toward these systems.

Trust is widely recognized as a critical prerequisite for successful human–AI interaction and technology adoption. In the context of AI, trust influences users' willingness to rely on AI recommendations, disclose personal information, and integrate AI systems into their daily activities (Glikson & Woolley, 2020). However, the complexity and opacity of many AI systems often create uncertainty regarding reliability, fairness, and accountability, potentially limiting public trust.

Consequently, researchers increasingly view trust in AI not only as a technical issue but also as a social and institutional phenomenon influenced by governance mechanisms, ethical standards, and regulatory frameworks (Scharowski et al., 2024).

2.2 The European Union Artificial Intelligence Act

The European Union Artificial Intelligence Act (EU AI Act) represents the first comprehensive legal framework specifically designed to regulate artificial intelligence systems (European Commission, 2024). The regulation adopts a risk-based approach that imposes different obligations depending on the level of risk associated with AI applications.

In addition to compliance requirements for developers and providers, the EU AI Act introduces transparency obligations intended to inform users when they are interacting with AI systems. These disclosures are designed to reduce information asymmetry and increase confidence in AI technologies by signalling compliance with recognized legal and ethical standards (European Commission, 2024).

However, while policymakers expect regulatory disclosures to increase trust, such information may also increase awareness of potential risks associated with AI technologies. Understanding

how users interpret regulatory disclosures therefore represents an important issue in contemporary AI governance research.

2.3 Perceived Risk in Artificial Intelligence

Perceived risk refers to an individual's subjective assessment of uncertainty and potential negative consequences associated with a particular technology or decision (Slovic, 1987). Risk Perception Theory suggests that individuals evaluate risks not only through objective information but also through cognitive and emotional processes that shape behavioural responses.

In the context of artificial intelligence, perceived risk has emerged as a major determinant of trust and technology acceptance. Previous studies have identified privacy concerns, algorithmic bias, lack of transparency, and cybersecurity threats as important sources of AI-related risk perceptions (Leschanowsky et al., 2024).

Research consistently demonstrates that higher perceived risk reduces users' trust in AI systems and their willingness to adopt AI technologies (Glikson & Woolley, 2020). Consequently, understanding and managing perceived risk has become an important objective for both researchers and policymakers.

2.4 Adoption Intention

Adoption intention refers to an individual's willingness or intention to use a particular technology in the future (Davis, 1989). Within the context of artificial intelligence, adoption intention reflects users' readiness to integrate AI systems into their daily activities, rely on AI-generated recommendations, and utilize AI-supported decision-making tools.

Previous studies have consistently identified trust as one of the most important determinants of AI adoption, particularly for technologies characterized by uncertainty and perceived risk (Glikson & Woolley, 2020). Consequently, the present study examines adoption intention as the final behavioural outcome of trust in AI systems.

2.5 Theoretical Foundations

This study draws upon four complementary theoretical perspectives to explain how EU AI Act disclosure influences perceived risk, trust in AI, and adoption intention. Specifically, the study integrates Risk Perception Theory, Signalling Theory, Trust Transfer Theory, and the

Technology Acceptance Model to explain the psychological and behavioural mechanisms underlying users' responses to AI regulation.

2.5.1 Risk Perception Theory

Risk Perception Theory explains how individuals subjectively evaluate uncertain or potentially harmful situations (Slovic, 1987). Unlike objective risk assessments based on statistical probabilities, perceived risk reflects psychological interpretations influenced by emotions, prior experiences, and available information. Technologies perceived as unfamiliar, complex, or difficult to understand are often associated with higher levels of perceived risk.

Artificial intelligence represents a particularly relevant context for Risk Perception Theory because AI systems often operate as complex and opaque "black boxes." Previous studies have identified privacy concerns, algorithmic bias, and lack of transparency as major sources of AI-related risk perceptions (Leschanowsky et al., 2024). Accordingly, this theory provides the foundation for examining the mediating role of perceived risk between EU AI Act disclosure and trust in AI systems.

2.5.2 Signalling Theory

Signalling Theory proposes that individuals rely on signals to reduce information asymmetry and uncertainty in decision-making situations (Spence, 1973). Signals such as certifications, labels, and regulatory approvals can influence perceptions regarding quality, reliability, and trustworthiness.

In the context of artificial intelligence, disclosures associated with the EU AI Act may function as signals of safety, transparency, and accountability that influence users' evaluations of AI systems. However, such disclosures may also increase awareness of potential risks associated with AI technologies. Therefore, Signalling Theory provides an important theoretical explanation for the potential effects of EU AI Act disclosure on trust and perceived risk.

2.5.3 Trust Transfer Theory

Trust Transfer Theory suggests that trust established in one institution or entity can be transferred to another related entity through perceived associations (Stewart, 2003). The theory has frequently been applied in digital services and electronic commerce to explain how trust in institutions influences trust in technologies and online platforms.

Within the context of this study, users who trust European institutions and regulatory authorities may transfer that trust to AI systems presented as compliant with the EU AI Act. Consequently, Trust Transfer Theory provides a theoretical basis for understanding how regulatory disclosure may influence trust in AI systems through institutional legitimacy.

2.5.4 Technology Acceptance Model

The Technology Acceptance Model (TAM) is one of the most widely used frameworks for explaining individuals' intentions to adopt new technologies (Davis, 1989). Although the original model focuses on perceived usefulness and perceived ease of use, later extensions have incorporated additional factors such as trust and perceived risk.

In the context of artificial intelligence, trust has been identified as an important determinant of technology adoption, particularly for systems characterized by uncertainty and limited transparency (Glikson & Woolley, 2020). Therefore, TAM provides the theoretical foundation for the proposed relationship between trust in AI and adoption intention in the present study.

2.6 Previous Empirical Evidence

Previous studies have shown that regulatory mechanisms, certification labels, and institutional safeguards can significantly influence consumer trust and risk perceptions by reducing uncertainty and increasing confidence in technologies and digital services. For example, Fox et al. (2022) found that General Data Protection Regulation (GDPR) labels improved users' trust in online services by signalling stronger privacy protection and institutional accountability.

Similar findings have been reported in the context of artificial intelligence. Scharowski et al. (2023) demonstrated that certification labels for trustworthy AI increased users' trust and willingness to use AI systems, suggesting that regulatory disclosures may function as signals of safety and reliability. Furthermore, Leschanowsky et al. (2024) identified privacy concerns, security issues, and transparency as major determinants of trust and adoption in AI technologies. Their findings indicate that higher levels of perceived risk are generally associated with lower trust and reduced willingness to adopt AI systems.

Despite these contributions, empirical evidence regarding AI-specific regulations remains limited. Existing studies have primarily focused on privacy labels, certification schemes, and

general trust in AI rather than legally mandated regulatory disclosures such as those introduced by the European Union Artificial Intelligence Act. Moreover, the role of perceived risk as a mediating mechanism between regulatory disclosure and trust remains underexplored, particularly among Generation Z users who represent one of the most active groups of AI users. Therefore, the present study experimentally examines the effects of EU AI Act disclosure on perceived risk, trust in AI, and adoption intention while investigating the mediating role of perceived risk in this relationship.

2.7 Hypotheses Development

Based on the theoretical foundations and previous empirical findings discussed in the literature review, this study proposes that EU AI Act disclosure may influence users' perceived risk and trust in AI systems. Furthermore, trust is expected to affect adoption intention, while perceived risk is proposed to mediate the relationship between regulatory disclosure and trust. Accordingly, the following hypotheses are developed:

- **H1:** EU AI Act disclosure is associated with trust in AI systems.
- **H2:** EU AI Act disclosure is associated with perceived risk of AI systems.
- **H3:** Perceived risk is associated with trust in AI systems.
- **H4:** Trust in AI is associated with adoption intention.
- **H5:** Perceived risk mediates the relationship between EU AI Act disclosure and trust in AI systems.

Chapter 3 : Research Methodology

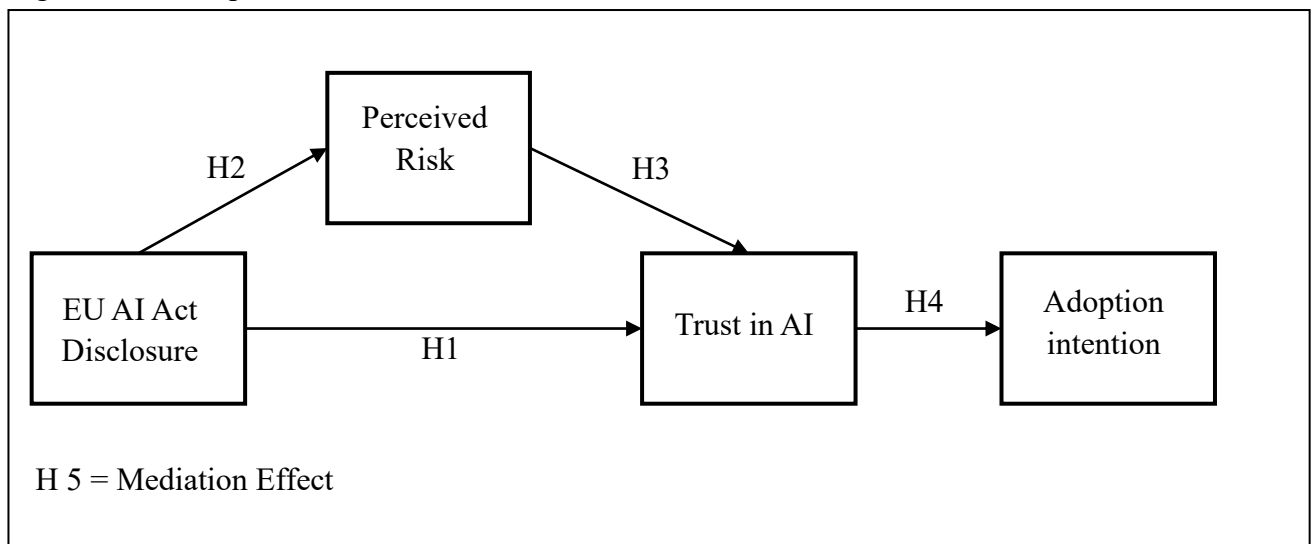
3.1 Research Design

This study adopted a quantitative experimental research design to examine the effects of EU AI Act disclosure on perceived risk, trust in AI, and adoption intention. A between-subjects experiment was conducted in which participants were randomly assigned to either a control group or a treatment group. Participants in the control group viewed a standard AI chatbot interface without regulatory information, whereas participants in the treatment group viewed the same chatbot interface accompanied by an EU AI Act disclosure statement. This design enabled the study to investigate the causal effects of regulatory disclosure on users' perceptions and behavioural intentions toward AI systems.

3.2 Conceptual Framework

This study proposes a conceptual framework to examine the relationships among EU AI Act disclosure, perceived risk, trust in AI, and adoption intention. Specifically, the framework assumes that EU AI Act disclosure influences users' perceived risk and trust in AI systems, while trust further affects adoption intention. Additionally, perceived risk is expected to mediate the relationship between EU AI Act disclosure and trust in AI systems. The conceptual framework of the study is presented in Figure 3.1.

Figure 3.1 Conceptual framework



Source: Developed by the author

3.3 Experimental Stimuli and Procedure

The study employed two experimental stimuli consisting of screenshots of an AI chatbot interface. Participants were randomly assigned to one of two experimental conditions. The

control group viewed a standard chatbot interface (see Appendix II) without any regulatory information, whereas the treatment group viewed the same chatbot interface (Appendix III) accompanied by an EU AI Act disclosure statement indicating compliance with European AI regulations. After viewing the assigned stimulus, participants completed an online questionnaire measuring perceived risk, trust in AI, adoption intention, and demographic characteristics. The survey was administered using the Qualtrics platform, which ensured standardized data collection and random assignment of participants to experimental conditions.

3.4 Sample and Data Collection

The target population of this study consisted of Generation Z individuals due to their high exposure to digital technologies and frequent use of AI applications. Data were collected through an online questionnaire (see Appendix I) distributed via the Qualtrics platform. A total of 30 participants completed the survey and were randomly assigned to one of the two experimental conditions, resulting in 15 participants in the control group and 15 participants in the EU AI Act disclosure group. The use of random assignment helped ensure comparability between the groups and increased the internal validity of the study.

3.5 Measurement and Operationalization of Variables

The study included four main constructs: EU AI Act disclosure, perceived risk, trust in AI, and adoption intention. EU AI Act disclosure was operationalized as the experimental manipulation, where participants were exposed either to a standard chatbot interface without regulatory information (control group) or to a chatbot interface containing EU AI Act disclosure information (treatment group). Perceived risk was measured using nine items covering privacy, fairness, and transparency concerns associated with AI systems.

Trust in AI was measured using six items assessing the reliability, confidence, and trustworthiness of the AI system, while adoption intention was measured using four items capturing participants' willingness to use AI technologies in the future. All measurement items were adapted from established scales in the literature and assessed using seven-point Likert scale. The operationalization of the study variables is summarized in Table 3.1.

Table 3.1 Operationalization of Variables

Variable	Role	Measurement
EU AI Act Disclosure	Independent Variable	Experimental manipulation (Control = 0, Treatment = 1)
Perceived Risk	Mediator Variable	9-item Likert scale (PR1–PR9)
Trust in AI	Dependent Variable and Predictor Variable	6-item Likert scale (TR1–TR6)
Adoption Intention	Dependent Variable	4-item Likert scale (AI1–AI4)

Source: Developed by the author

3.6 Data Analysis Methods

The collected data were analysed using descriptive and inferential statistical techniques. Descriptive statistics were used to summarize participants' demographic characteristics and the distribution of the study variables. Reliability analysis using Cronbach's alpha was conducted to assess the internal consistency of the measurement scales. Hypotheses were tested using independent sample t-tests, linear regression analysis, and mediation analysis. All statistical analyses were performed using the R statistical software environment.

Chapter 4 : Results and Findings

4.1 Demographic Characteristics

The final sample consisted of 30 Generation Z participants who were equally distributed between the control condition and the EU AI Act disclosure condition, with 15 participants assigned to each group. The participants had a mean age of 23.03 years ($SD = 2.92$), indicating that the sample primarily represented young adults with high exposure to digital technologies. Regarding gender distribution, 66.7% of participants were male, 26.7% were female, and 6.7% identified as other. Participants also reported relatively high levels of prior experience with AI technologies ($M = 5.47$, $SD = 1.17$), suggesting familiarity with AI applications and tools. A summary of the demographic characteristics of the sample is presented in Table 4.1.

Table 4.1 Demographic Characteristics of the Sample

Variable	Category	n	%
Gender	Male	20	66.7
	Female	8	26.7
	Other	2	6.7
Experimental Condition	Control	15	50.0
	EU AI Act Disclosure	15	50.0
Age	Mean (SD)	23.03 (2.92)	
AI Experience	Mean (SD)	5.47 (1.17)	

Source: Author's own calculations based on the survey data (2026).

4.2 Reliability Analysis and Descriptive Statistics

Prior to hypothesis testing, the reliability and descriptive statistics of the study variables were examined. Reliability was assessed using Cronbach's alpha coefficients to evaluate the internal consistency of the measurement scales. All constructs demonstrated high levels of reliability, with Cronbach's alpha values ranging from 0.90 to 0.95, exceeding the recommended threshold of 0.70 (Nunnally & Bernstein, 1994). Regarding descriptive statistics, participants reported moderate levels of perceived risk, trust in AI, and adoption intention. These results indicate that

the measurement scales were reliable and suitable for subsequent statistical analyses. Table 4.2 presents the reliability and descriptive statistics of the study variables.

Table 4.2 Reliability Analysis and Descriptive Statistics

<i>Construct</i>	<i>Items</i>	<i>Cronbach's α</i>	<i>Mean</i>	<i>SD</i>
<i>Perceived Risk</i>	9	0.94	4.04	0.93
<i>Trust in AI</i>	6	0.91	3.59	0.85
<i>Adoption Intention</i>	4	0.91	4.07	0.86

Source: Author's own calculations based on the survey data (2026).

4.3 Correlation Analysis

Pearson correlation analysis was conducted to examine the relationships among the main study variables. As presented in Table 4.3, perceived risk was significantly and negatively correlated with trust in AI ($r = -0.550$, $p = .002$), indicating that higher levels of perceived risk were associated with lower levels of trust in AI systems. Similarly, perceived risk exhibited a significant negative relationship with adoption intention ($r = -0.445$, $p = .014$). In contrast, trust in AI demonstrated a strong positive correlation with adoption intention ($r = .683$, $p < .001$), suggesting that individuals with higher levels of trust in AI were more willing to adopt AI technologies. These findings provide preliminary support for the proposed relationships and justify the subsequent hypothesis testing analyses.

Table 4.3 Correlation Matrix

<i>Variable</i>	<i>1</i>	<i>2</i>	<i>3</i>
<i>1. Perceived Risk</i>	—		
<i>2. Trust in AI</i>	-0.550**	—	
<i>3. Adoption Intention</i>	-0.445*	0.683***	—

Source: Author's own calculations based on the survey data (2026).

Note: * $p < .05$, ** $p < .01$, *** $p < .001$

4.4 Independent samples t-test of EU AI Act Disclosure and Trust in AI

An independent samples t-test was conducted to examine whether EU AI Act disclosure influenced trust in AI systems. The results revealed a statistically significant difference in trust

between the control group and the EU AI Act disclosure group, $t(28) = 2.41$, $p = .023$. Participants in the control group reported higher levels of trust in AI ($M = 3.94$, $SD = 0.66$) compared to participants exposed to the EU AI Act disclosure ($M = 3.24$, $SD = 0.91$).

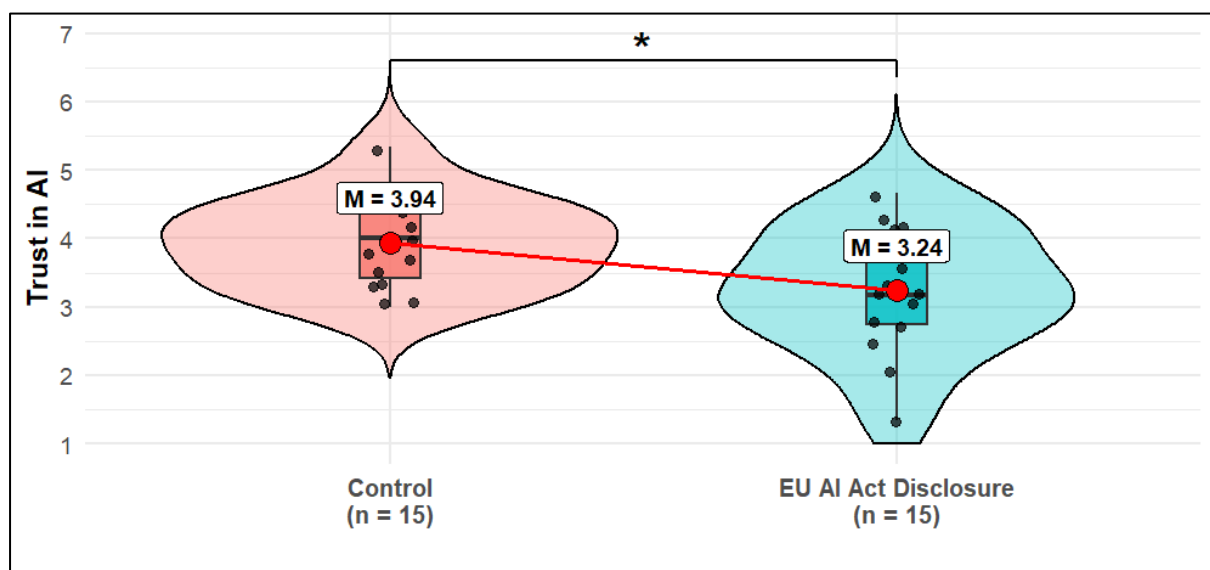
Table 4.4 Independent Samples t-test for H1: EU AI Act Disclosure and Trust in AI

<i>Group</i>	<i>N</i>	<i>Mean</i>	<i>SD</i>	<i>t(df)</i>	<i>p</i>	<i>95% CI</i>
<i>Control</i>	15	3.94	0.66	-	-	-
<i>EU AI Act Disclosure</i>	15	3.24	0.91	2.41 (28)	.023	[0.10, 1.30]

Source: Author's own calculations based on the survey data (2026).

Figure 4.1 illustrates the differences in trust in AI between the control group and the EU AI Act disclosure group. The difference was statistically significant, $t(28) = 2.41$, $p = .023$, also indicating that exposure to EU AI Act information reduced trust in AI systems.

Figure 4.1 Trust in AI by Experimental Condition



Source: Author's own calculations based on the survey data (2026).

Note: * $p < .05$, ** $p < .01$, *** $p < .001$

4.5 Independent samples t-test of EU AI Act Disclosure and Perceived Risk

An independent samples t-test was conducted to examine whether EU AI Act disclosure influenced perceived risk associated with AI systems. The results indicated a statistically significant difference in perceived risk between the control group and the EU AI Act disclosure group, $t(28) = -2.52$, $p = .018$. Participants exposed to the EU AI Act disclosure reported

significantly higher levels of perceived risk ($M = 4.44$, $SD = 0.81$) compared to participants in the control condition ($M = 3.65$, $SD = 0.90$).

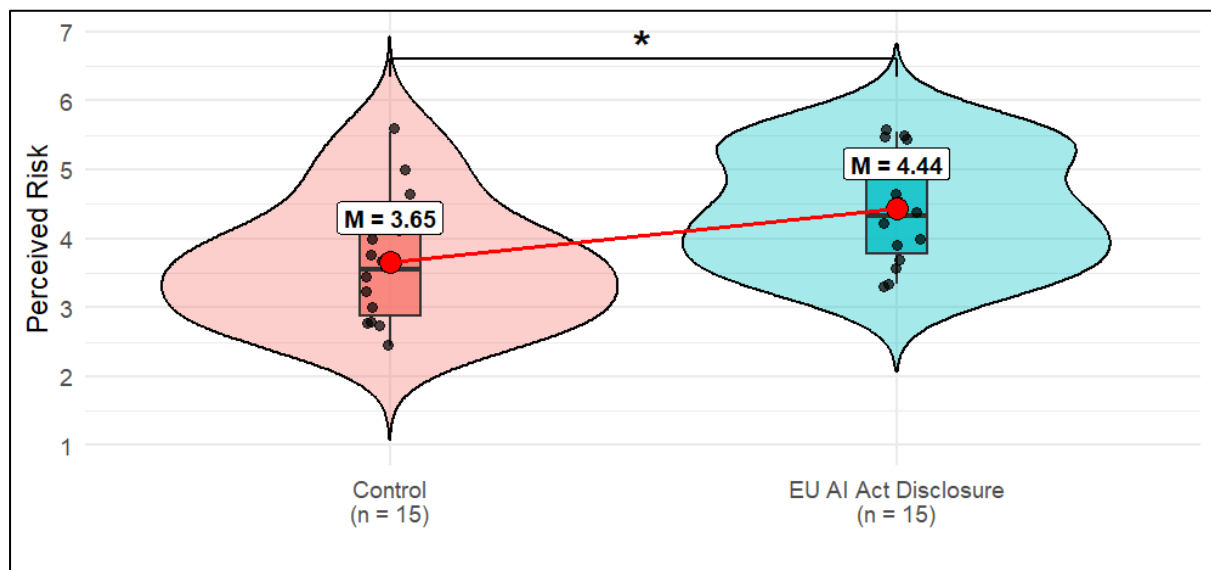
Table 4.5 Independent Samples t-test for H2: EU AI Act Disclosure and Perceived Risk

<i>Group</i>	<i>N</i>	<i>Mean</i>	<i>SD</i>	<i>t(df)</i>	<i>p</i>	<i>95% CI</i>
<i>Control</i>	15	3.65	0.90			
<i>EU AI Act Disclosure</i>	15	4.44	0.81	-2.52 (28)	.018	[-1.42, -0.15]

Source: Author's own calculations based on the survey data (2026).

Figure 4.2 illustrates the differences in perceived risk between the control group and the EU AI Act disclosure group. The difference was statistically significant, $t(28) = -2.52$, $p = .018$, also indicating that exposure to EU AI Act information increased perceived risk associated with AI systems.

Figure 4.2 Perceived Risk by Experimental Condition



Source: Author's own calculations based on the survey data (2026).

Note: * $p < .05$, ** $p < .01$, *** $p < .001$

4.6 Regression analysis of Perceived Risk and Trust in AI

A simple linear regression analysis was conducted to examine the relationship between perceived risk and trust in AI systems. The results indicated that perceived risk had a significant negative effect on trust in AI ($B = -0.509$, $SE = 0.146$, $\beta = -0.550$, $t(28) = -3.49$, $p = .002$). The model explained approximately 30.3% of the variance in trust in AI ($R^2 = .303$, Adjusted $R^2 = .278$), and the overall regression model was statistically significant, $F(1, 28) = 12.15$, $p = .002$.

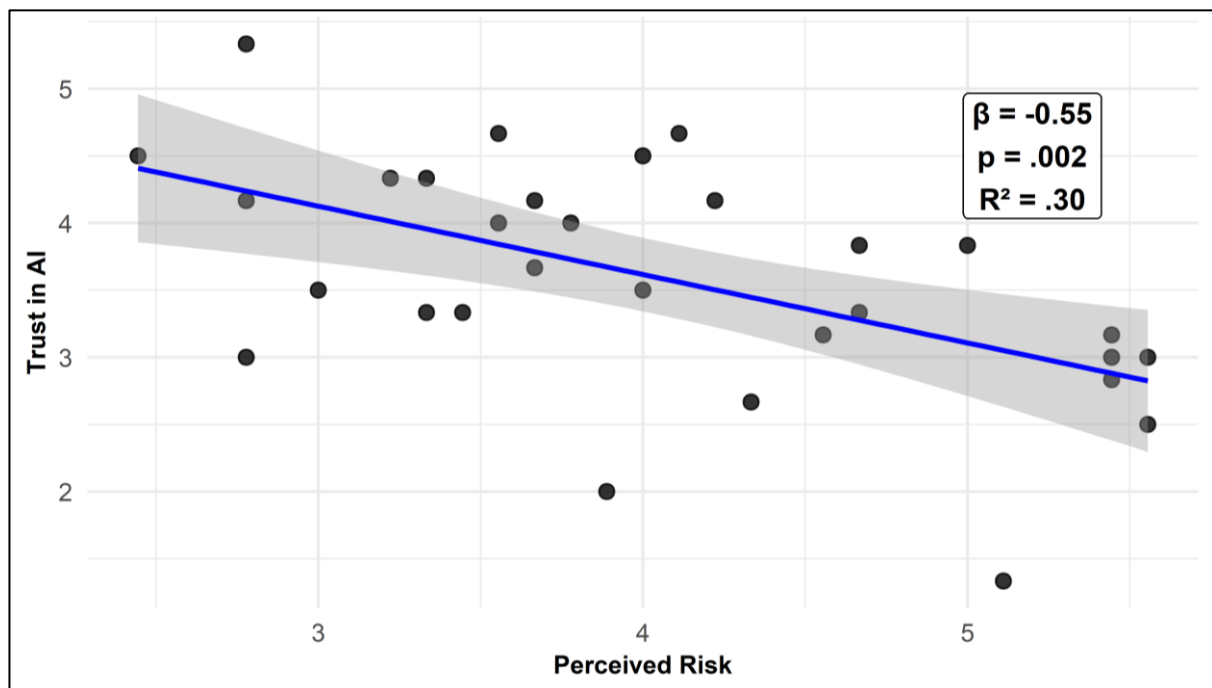
Table 4.6 Regression Results for H3: Perceived Risk and Trust in AI

<i>Predictor</i>	<i>B</i>	<i>SE</i>	β	<i>t</i>	<i>p</i>
<i>Constant</i>	5.653	0.606	—	9.34	< .001
<i>Perceived Risk</i>	-0.509	0.146	-0.550	-3.49	.002
<i>Model Statistics</i>	$R^2=.303$	Adjusted $R^2=.278$	$F(1, 28) = 12.15$		$p = .002$

Source: Author's own calculations based on the survey data (2026).

Figure 4.3 illustrates the relationship between perceived risk and trust in AI. The regression line indicates a negative association between the two variables, suggesting that participants who perceived higher levels of risk reported lower levels of trust in AI systems. The model explained approximately 30.3% of the variance in trust in AI ($R^2 = .30$).

Figure 4.3 Effect of Perceived Risk on Trust in AI



Source: Author's own calculations based on the survey data (2026).

4.7 Regression analysis of Trust in AI and Adoption Intention

A simple linear regression analysis was conducted to examine the relationship between trust in AI and adoption intention. The results revealed that trust in AI had a significant positive effect on adoption intention ($B = 0.680$, $SE = 0.138$, $\beta = 0.683$, $t(28) = 4.94$, $p < .001$). The model explained approximately 46.6% of the variance in adoption intention ($R^2 = .466$, Adjusted $R^2 = .447$), and the overall regression model was statistically significant, $F(1, 28) = 24.42$, $p < .001$.

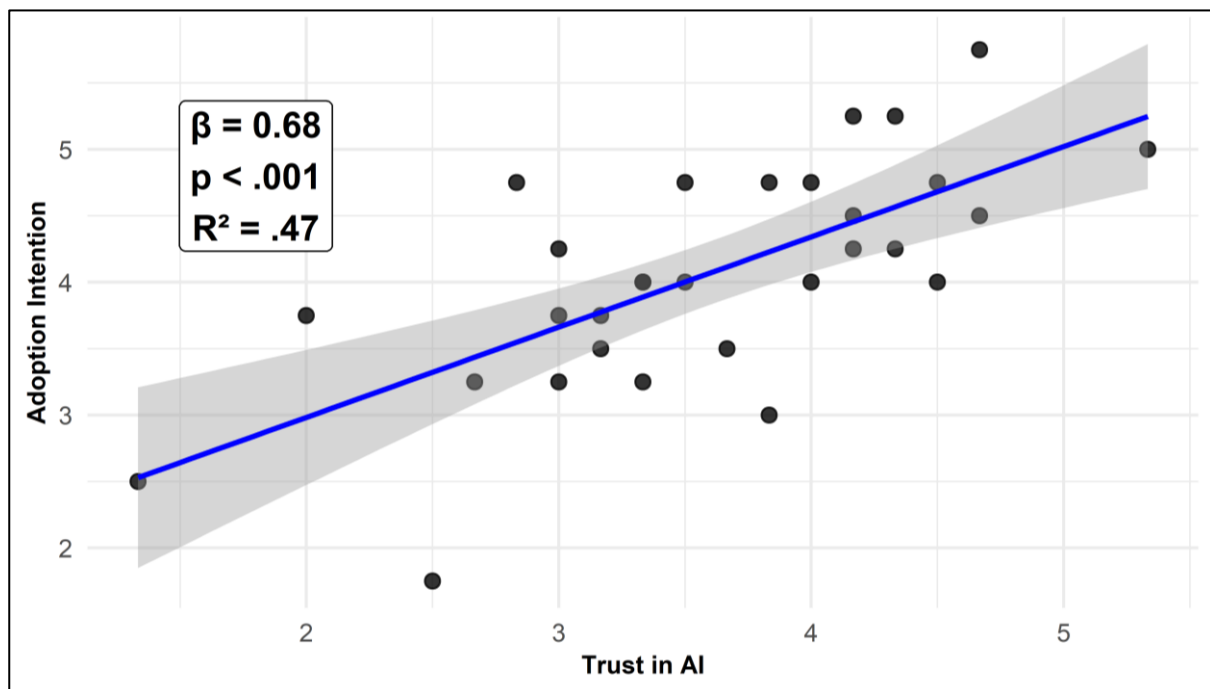
Table 4.7 Regression Results for H4: Trust in AI and Adoption Intention

<i>Predictor</i>	<i>B</i>	<i>SE</i>	β	<i>t</i>	<i>p</i>
<i>Constant</i>	1.624	0.508	—	3.20	.003
<i>Trust in AI</i>	0.680	0.138	0.683	4.94	< .001
<i>Model Statistics</i>	$R^2 = .466$	Adjusted $R^2 = .447$	$F(1, 28) = 24.42$		$p < .001$

Source: Author's own calculations based on the survey data (2026).

Figure 4.4 illustrates the relationship between trust in AI and adoption intention. The regression line indicates a positive association between the two variables, suggesting that participants who reported higher levels of trust in AI also reported stronger intentions to adopt and use AI systems. The model explained approximately 46.6% of the variance in adoption intention ($R^2 = .47$).

Figure 4.4 Effect of Trust in AI on Adoption Intention



Source: Author's own calculations based on the survey data (2026).

4.8 Mediation analysis of Perceived Risk

A mediation analysis using nonparametric bootstrap estimation with 5,000 resamples was conducted to examine whether perceived risk mediated the relationship between EU AI Act disclosure and trust in AI systems. The results revealed a significant indirect effect of EU AI Act disclosure on trust through perceived risk (ACME = -0.332, 95% CI [-0.748, -0.046], $p = .020$). The direct effect of EU AI Act disclosure on trust was not statistically significant after

controlling for perceived risk (ADE = -0.368, 95% CI [-0.959, 0.180], $p = .166$). The total effect of EU AI Act disclosure on trust remained significant (Total Effect = -0.700, 95% CI [-1.283, -0.164], $p = .012$).

The mediation analysis further indicated that perceived risk accounted for approximately 47.4% of the total effect of EU AI Act disclosure on trust in AI systems (Proportion Mediated = 47.4%, $p = .029$).

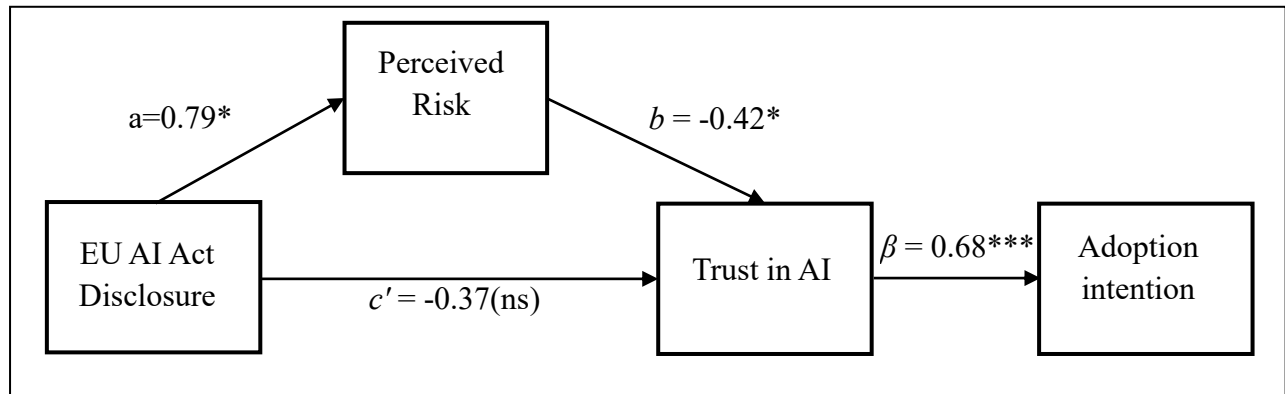
Table 4.8 Mediation Analysis Results for H5: The Mediating Role of Perceived Risk

<i>Effect</i>	<i>Estimate</i>	<i>95% CI Lower</i>	<i>95% CI Upper</i>	<i>p</i>
<i>ACME (Indirect Effect)</i>	-0.332	-0.748	-0.046	.020
<i>ADE (Direct Effect)</i>	-0.368	-0.959	0.180	.166
<i>Total Effect</i>	-0.700	-1.283	-0.164	.012
<i>Proportion Mediated</i>	0.474	0.066	1.581	.029

Source: Author's own calculations based on the survey data (2026).

Figure 4.5 illustrates the mediating role of perceived risk in the relationship between EU AI Act disclosure and trust in AI systems. Exposure to the EU AI Act disclosure significantly increased perceived risk ($a = 0.79$, $p = .018$). In turn, perceived risk significantly reduced trust in AI systems ($b = -0.42$, $p = .014$). Although the total effect of EU AI Act disclosure on trust was significant ($c = -0.70$, $p = .023$), the direct effect became statistically non-significant after controlling for perceived risk ($c' = -0.37$, $p = .218$). These findings indicate that perceived risk fully mediates the relationship between EU AI Act disclosure and trust in AI systems. Furthermore, trust in AI had a significant positive effect on adoption intention ($\beta = 0.68$, $p < .001$).

Figure 4.5 Mediation Model of the Relationship Between EU AI Act Disclosure, Perceived Risk, Trust in AI, and Adoption Intention



Source: Author's own calculations based on the survey data (2026).

Note: $p < .05$ *, $p < .01$ **, $p < .001$ ***; ns = not significant.

Chapter 5 : Discussion of Findings

5.1. EU AI Act Disclosure, Trust, and Perceived Risk

The findings revealed that exposure to EU AI Act disclosure significantly reduced trust in AI systems while simultaneously increasing perceived risk among participants. These findings are particularly interesting because regulatory disclosures are generally expected to reassure users and increase confidence in emerging technologies. Instead, the results suggest that providing information about AI regulation may unintentionally increase awareness of potential risks associated with artificial intelligence, including privacy concerns, algorithmic bias, and lack of transparency.

This finding partially contradicts previous studies examining certification labels and privacy regulations, which generally reported positive effects on trust and technology acceptance (Fox et al., 2022; Scharowski et al., 2023). However, the results are consistent with Risk Perception Theory, which argues that additional information regarding uncertainty and potential harm may increase perceived risk rather than reduce it (Slovic, 1987). The findings therefore suggest that regulatory disclosures may operate not only as signals of safety but also as reminders of the risks that motivated regulation in the first place.

5.2 The Relationship Between Perceived Risk and Trust in AI

Consistent with previous research, perceived risk was found to have a significant negative effect on trust in AI systems. Participants who perceived higher levels of risk reported lower levels of trust in AI technologies. This finding supports previous studies highlighting privacy concerns, security risks, and transparency issues as major determinants of trust in AI systems (Leschanowsky et al., 2024).

The results also provide empirical support for Risk Perception Theory by demonstrating that risk perceptions play an important role in shaping trust formation in emerging technologies. In the context of AI, users often possess limited understanding of how algorithms generate recommendations and decisions, making risk perceptions particularly influential in evaluating the trustworthiness of AI systems.

5.3 The Relationship Between Trust and Adoption Intention

The findings further demonstrated that trust in AI significantly and positively influenced adoption intention. Participants who reported higher levels of trust were more willing to adopt

and use AI technologies in the future. This result is consistent with previous extensions of the Technology Acceptance Model, which identify trust as an important determinant of technology adoption in uncertain environments (Davis, 1989; Glikson & Woolley, 2020).

The findings suggest that trust acts as an important prerequisite for public acceptance of artificial intelligence. Even highly capable AI systems may face resistance if users do not perceive them as reliable, transparent, and trustworthy.

5.4 The Mediating Role of Perceived Risk

One of the most important findings of this study was the mediating role of perceived risk in the relationship between EU AI Act disclosure and trust in AI systems. The mediation analysis revealed that the direct effect of EU AI Act disclosure on trust became statistically non-significant after accounting for perceived risk, indicating that the effect of disclosure operated primarily through changes in risk perceptions.

This finding suggests that regulatory disclosures do not directly reduce trust. Instead, disclosures influence trust indirectly by shaping users' perceptions of AI-related risks. Consequently, policymakers should recognize that transparency alone may not be sufficient to increase public trust unless accompanied by communication strategies that also address and mitigate users' concerns regarding AI technologies.

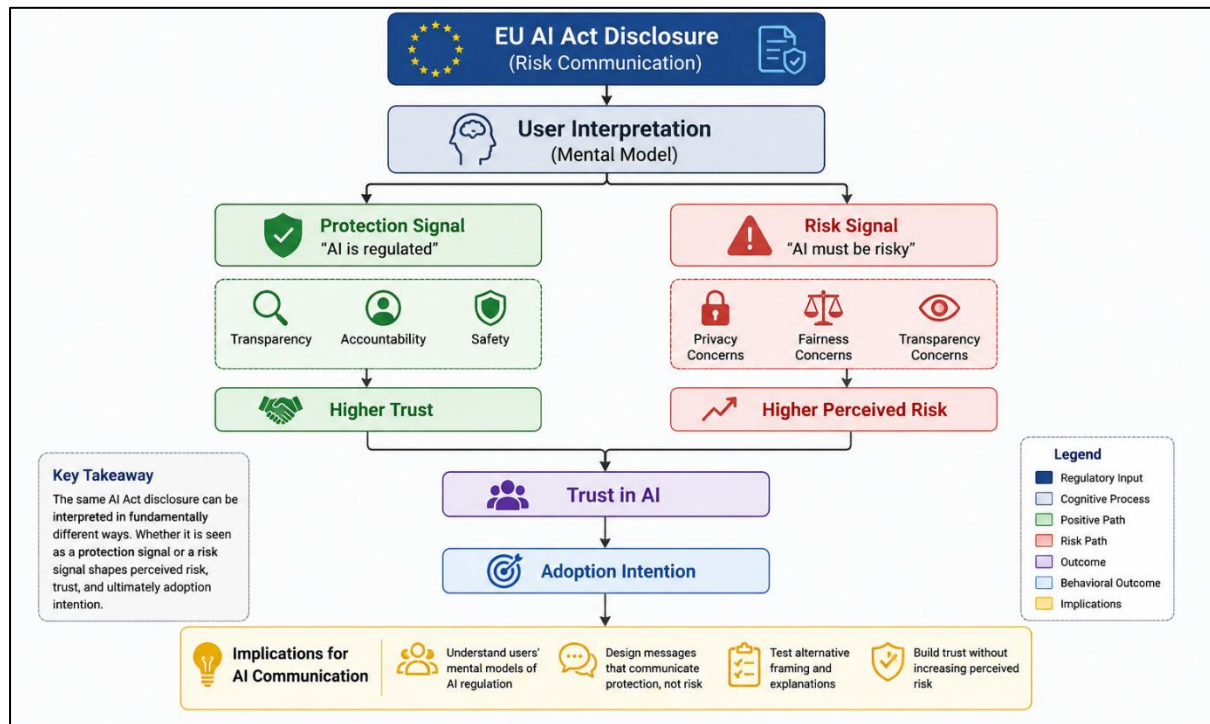
5.5 Theoretical Implications

This study contributes to the emerging literature on AI governance by providing empirical evidence on how regulatory disclosures influence trust in AI systems through psychological mechanisms. While previous studies have generally argued that regulations and certifications increase trust and acceptance, the present findings demonstrate that regulatory disclosures may also unintentionally increase perceptions of risk.

The findings support Risk Perception Theory by showing that higher perceived risk significantly reduces trust in AI systems. In addition, the results extend Signalling Theory by suggesting that regulatory disclosures do not always operate as positive signals of safety and reliability. Instead, participants may interpret the disclosure itself as an indication that significant risks exist and therefore require regulatory intervention.

Furthermore, the study contributes to Trust Transfer Theory and the Technology Acceptance Model by highlighting the importance of trust in shaping adoption intentions. By integrating regulatory disclosure, perceived risk, trust, and adoption intention into a single framework, the study provides a more comprehensive understanding of public responses to AI regulation.

Figure 5.1 User Interpretation Mechanism of EU AI Act Disclosure as Protection and Risk Signals.



Source: Developed by the author using AI tools (ChatGPT)

An additional contribution of this study is the proposed interpretation mechanism illustrated in Figure 5.1. Although the EU AI Act disclosure was designed to function as a protection signal communicating transparency, accountability, and safety, participants may have interpreted the disclosure as a risk signal indicating the existence of privacy concerns, algorithmic bias, and transparency issues. This cognitive interpretation process suggests that users do not respond directly to regulatory information itself but rather to the mental models and meanings they attach to such information. Consequently, the effects of AI regulation depend not only on regulatory design but also on how regulatory messages are framed, communicated, and interpreted by users.

5.6 Practical Implications

The findings generate important implications for policymakers, regulators, and AI developers involved in the implementation of the EU AI Act. Although the regulation was designed to

function as a protection mechanism that increases public confidence in AI systems, participants may have interpreted the disclosure as a signal of risk rather than safety.

This finding suggests that the effectiveness of AI regulation depends not only on the regulation itself but also on how regulatory information is communicated and interpreted by users. Disclosure alone may therefore be insufficient to increase trust and public acceptance of AI technologies.

For AI developers and policymakers, the results highlight the importance of complementing transparency disclosures with clear explanations regarding safety measures, fairness safeguards, privacy protections, and human oversight mechanisms. Effective communication strategies may therefore play an essential role in ensuring that AI regulation promotes trust rather than unintentionally increasing concerns regarding AI systems.

Chapter 6 : Conclusion

6.1 Conclusion

This study examined the effects of EU AI Act disclosure on perceived risk, trust in AI, and adoption intention among Generation Z users. The findings revealed that exposure to EU AI Act disclosure significantly increased perceived risk while reducing trust in AI systems. Furthermore, perceived risk negatively influenced trust, whereas trust positively affected adoption intention.

The mediation analysis demonstrated that perceived risk acted as an important psychological mechanism linking regulatory disclosure and trust in AI systems. Although the EU AI Act was intended to increase confidence in AI technologies through transparency and accountability, participants appeared to interpret the disclosure as an indication that AI systems involve substantial risks that require regulation. Consequently, the disclosure functioned more as a risk signal than as a protection signal, leading to higher perceived risk and lower trust. Overall, the empirical findings supported all five hypotheses, confirming the proposed relationships among EU AI Act disclosure, perceived risk, trust in AI, and adoption intention.

The findings suggest that the effectiveness of AI regulation depends not only on the existence of regulatory frameworks but also on how regulatory information is communicated and interpreted by users. Regulatory transparency alone may therefore be insufficient to promote public trust unless accompanied by communication strategies that emphasize safety measures, privacy protections, fairness safeguards, and human oversight mechanisms.

6.2 Limitations and Future Research

Several limitations of this study provide opportunities for future research. First, the study relied on a relatively small sample consisting exclusively of Generation Z participants, which may limit the generalizability of the findings to other demographic groups. Future studies should therefore examine whether older generations and individuals with different levels of technological experience respond similarly to AI regulatory disclosures.

Second, the study employed a single experimental scenario involving one specific form of EU AI Act disclosure. Different disclosure formats, message framing strategies, and communication approaches may produce different effects on perceived risk and trust in AI

systems. Future research could investigate how variations in the presentation and wording of regulatory information influence public responses to AI governance initiatives.

Finally, future studies should further explore the cognitive interpretation processes underlying users' responses to AI regulation. The findings of the present study suggest that users may interpret regulatory disclosures either as protection signals that increase trust or as risk signals that increase concern and uncertainty. Understanding the factors that determine these interpretations, including AI literacy, prior experience with AI, and institutional trust, may provide important insights for the effective communication of AI regulations and policies.

Declaration on the Use of AI

I hereby declare that, in the planning, drafting, and revision of this master's thesis, I made use of the generative AI tool ChatGPT for the following purposes:

- To identify relevant academic sources and literature, including DOIs.
- To assist in planning the structure and overall organization of the thesis.
- To generate and refine R code for statistical data analysis.
- To generate visual illustrations, such as the User Interpretation Mechanism.
- To improve the language, grammar, and clarity of my own phrases, sentences, and paragraphs.

I hereby certify that this declaration is complete to the best of my knowledge. All work has been reviewed by me, and I take full responsibility for it. I understand that the omission or misrepresentation of my reliance on AI will constitute fraud.

Munich, 15.07.2026

(Place, Date)


(Signature)

References

- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340.
- European Commission. (2024). *Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act)*. Official Journal of the European Union.
- Fox, G., Lynn, T., & Rosati, P. (2022). Enhancing consumer perceptions of privacy and trust: A GDPR label perspective. *Information Technology & People*, 35(8), 181–204.
- Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of Management Annals*, 14(2), 627–660.
- Leschanowsky, A., Rech, S., Popp, B., & Bäckström, T. (2024). Evaluating privacy, security, and trust perceptions in conversational AI: A systematic review. *Computers in Human Behavior*, 159, 108344.
- Liu, Y., Wang, P. W., & Chen, Y. Y. (2023). Trust in government, perceived integrity and food safety protective behavior: The mediating role of risk perception. *International Journal of Public Health*, 68, 1605432.
- Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *Academy of Management Review*, 20(3), 709–734.
- Nunnally, J. C., & Bernstein, I. H. (1994). *Psychometric theory* (3rd ed.). McGraw-Hill
- Organisation for Economic Co-operation and Development. (2019). *Recommendation of the Council on Artificial Intelligence*. OECD Publishing.
- Scharowski, N., Benk, M., Kühne, S. J., Wettstein, L., & Brühlmann, F. (2023). Certification labels for trustworthy AI: Insights from an empirical mixed-method study. In *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency* (pp. 248–260). ACM.
- Scharowski, N., Perrig, S. A. C., Aeschbach, L. F., von Felten, N., Opwis, K., Wintersberger, P., & Brühlmann, F. (2024). To trust or distrust trust measures: Validating questionnaires for trust in AI. *arXiv*.
- Slovic, P. (1987). Perception of risk. *Science*, 236(4799), 280–285.

- Spence, M. (1973). Job market signaling. *The Quarterly Journal of Economics*, 87(3), 355–374.
- Stewart, K. J. (2003). Trust transfer on the World Wide Web. *Organization Science*, 14(1), 5–17.

Appendix

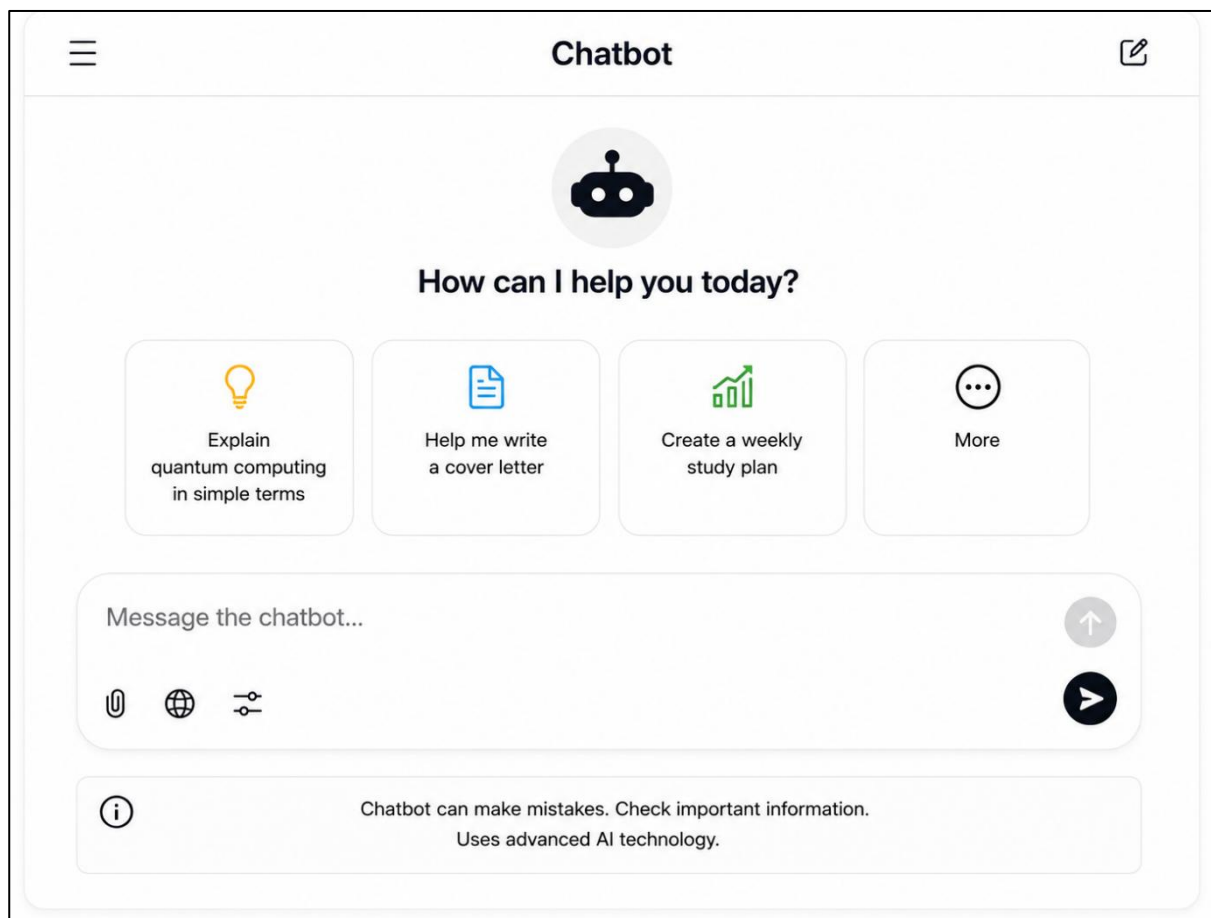
Appendix I : Survey Questionnaire

Section	Question / Item	Response Format
Participant Information and Consent	Participants were informed about the purpose of the study, voluntary participation, anonymity, confidentiality, and data protection procedures.	Information page
Experimental Manipulation	Participants were randomly assigned to one of two chatbot conditions: (1) AI chatbot with an EU AI Act compliance disclosure or (2) AI chatbot without such disclosure.	Random assignment
Demographic Information	What is your age?	Open numeric response
	What is your gender?	Male / Female / Other / Prefer not to say
AI Experience	How frequently do you use AI applications?	1 = Never to 7 = Very Frequently
Perceived Risk (PR1)	I am concerned that using this AI system could compromise my privacy.	7-point Likert scale
Perceived Risk (PR2)	I am worried that my personal data may be misused by this AI system.	7-point Likert scale
Perceived Risk (PR3)	I believe this AI system could make unfair decisions.	7-point Likert scale
Perceived Risk (PR4)	I am concerned that this AI system may contain biases.	7-point Likert scale
Perceived Risk (PR5)	I am uncertain about the reliability of this AI system.	7-point Likert scale
Perceived Risk (PR6)	I think using this AI system involves risks.	7-point Likert scale
Perceived Risk (PR7)	I find it difficult to understand how this AI system works.	7-point Likert scale
Perceived Risk (PR8)	The operation of this AI system lacks transparency.	7-point Likert scale

Perceived Risk (PR9)	Overall, I perceive this AI system as risky to use.	7-point Likert scale
Trust in AI (TR1)	I trust the information provided by this AI system.	7-point Likert scale
Trust in AI (TR2)	This AI system behaves reliably.	7-point Likert scale
Trust in AI (TR3)	I feel confident using this AI system.	7-point Likert scale
Trust in AI (TR4)	I consider this AI system trustworthy.	7-point Likert scale
Trust in AI (TR5)	This AI system acts in the user's best interest.	7-point Likert scale
Trust in AI (TR6)	I would feel comfortable relying on this AI system regularly.	7-point Likert scale
Adoption Intention (AI1)	I intend to use this AI system in the future.	7-point Likert scale
Adoption Intention (AI2)	I would consider using this AI system for important tasks.	7-point Likert scale
Adoption Intention (AI3)	I would recommend this AI system to others.	7-point Likert scale
Adoption Intention (AI4)	Overall, I am willing to adopt this AI system.	7-point Likert scale

Note. All attitudinal items were measured using a seven-point Likert scale ranging from **1 = Strongly disagree** to **7 = Strongly agree**.

Appendix II : Control group (standard chatbot)



Appendix III : Treatment group (EU AI Act disclosure)

Chatbot



This AI system operates under the European Union Artificial Intelligence Act (EU AI Act).

 Transparency

 Human oversight

 Accountability





How can I help you today?



Explain quantum computing in simple terms



Help me write a cover letter



Create a weekly study plan



More

Message the chatbot...









Chatbot can make mistakes. Check important information.
Uses advanced AI technology. |  **Operates under the EU AI Act.**

IV

Declaration of authorship

Last Name: Katupulle Gedara

First Name: Madushan Sanjeewa Bandara

Date of Birth: 11.03.1998

I hereby declare that the thesis submitted is my own unaided work. All direct or indirect sources used are acknowledged as references.

I am aware that the thesis, in digital form, can be examined for the use of unauthorized aid and artificial intelligence to determine whether the thesis as a whole or parts incorporated into it may be deemed plagiarism. For the comparison of my work with existing sources, I agree that it shall be entered in a database where it shall also remain after examination, to enable comparison with future theses submitted. Also, grant permission for the final version of this thesis to be accessed and used by the university for examination, grading, grade management, and other internal academic and administrative purposes in accordance with the applicable university regulations.

Further rights of reproduction and usage, however, are not granted here.

This paper was not previously presented to another examination board and has not been published.

Munich, 15.07.2026

(Place, Date)



(Signature)