

An Ethical Information System Success Framework for Government Artificial Intelligence

Dwi Yuniarto ^{(1)*}, A'ang Subiyakto ⁽²⁾, Aedah Abd. Rahman ⁽³⁾

¹ Department of Informatics, Sebelas April University, Sumedang, Indonesia

² Department of Information Systems, UIN Syarif Hidayatullah, Jakarta, Indonesia

³ Department of ICT, Asia E University, Selangor, Malaysia

e-mail : dwiyuniarto@unsap.ac.id, aang_subiyakto@uinjkt.ac.id,
aedah.abdrahman@aeu.edu.my.

* Corresponding author.

This article was submitted on 21 January 2026, revised on 11 June 2026, accepted on 17 June 2026, and published on 25 September 2026.

Abstract

The rapid adoption of Artificial Intelligence (AI) in government has transformed public service delivery and decision-making while raising ethical concerns related to fairness, transparency, accountability, and privacy. This study develops and validates an Ethical Information System Success Framework for Artificial Intelligence in Government by integrating Ethical AI governance principles with the Information System Success Model. A sequential mixed-method approach with triangulation-based instrument testing was employed. The model was developed using an Input–Process–Output logic and operationalized into 11 latent variables with 55 measurement items. Quantitative data were collected from 86 government employees using AI-based systems during October–November and analyzed using PLS-SEM, complemented by qualitative expert validation. The findings reveal that ethical transparency, accountability, and privacy and security significantly influence user trust, whereas ethical fairness does not demonstrate a statistically significant effect. User trust and satisfaction subsequently drive actual AI use, which strongly contributes to net public benefits. The empirical study was conducted among Indonesian government employees using AI-supported public service systems.

Keywords: Ethical Artificial Intelligence, Information System Success, Government, Trust, Digital Governance

Abstrak

Pesatnya adopsi Artificial Intelligence (AI) dalam pemerintahan telah mentransformasi penyelenggaraan layanan publik dan proses pengambilan keputusan, sekaligus memunculkan berbagai isu etis terkait keadilan, transparansi, akuntabilitas, dan privasi. Penelitian ini bertujuan mengembangkan dan memvalidasi Ethical Information System Success Framework for Artificial Intelligence in Government dengan mengintegrasikan prinsip tata kelola AI yang etis ke dalam model Information System Success. Penelitian menggunakan pendekatan mixed-method berurutan dengan pengujian instrumen berbasis triangulasi. Model dikembangkan menggunakan logika Input–Process–Output dan dioperasionalkan ke dalam 11 variabel laten dengan 55 butir pengukuran. Data kuantitatif dikumpulkan dari 86 pegawai pemerintah pengguna sistem berbasis AI selama periode Oktober–November dan dianalisis menggunakan PLS-SEM, yang dilengkapi dengan validasi kualitatif oleh para pakar. Hasil penelitian menunjukkan bahwa transparansi etis, akuntabilitas, serta privasi dan keamanan berpengaruh signifikan terhadap kepercayaan pengguna, sedangkan keadilan etis tidak menunjukkan pengaruh yang signifikan secara statistik. Kepercayaan dan kepuasan pengguna selanjutnya mendorong penggunaan aktual AI yang berkontribusi kuat terhadap manfaat bersih yang diperoleh organisasi. Studi empiris ini dilakukan pada pegawai pemerintah di Indonesia yang menggunakan sistem layanan publik berbasis AI.

Kata Kunci: Artificial Intelligence Etis, Keberhasilan Sistem Informasi, Pemerintahan, Kepercayaan, Tata Kelola Digital



1. INTRODUCTION

Artificial Intelligence (AI) has become a core digital infrastructure in contemporary government administration, supporting policy analysis, public service delivery, surveillance, forecasting, decision support systems, and smart governance initiatives. Governments worldwide increasingly rely on AI to enhance efficiency, accuracy, scalability, and responsiveness of public services (Al-Ansi et al., 2024; Al-Besher & Kumar, 2022; Anshari et al., 2025; Yigitcanlar et al., 2024). However, alongside its transformative potential, AI deployment in the public sector also introduces significant ethical, legal, and governance risks, including algorithmic bias, lack of transparency, data privacy violations, accountability ambiguity, and erosion of public trust (Al-Dulaimi & Mohammed, 2026; Vatamanu & Tofan, 2025). These challenges indicate that the success of AI in government cannot be evaluated merely through technical or performance-based indicators but must be assessed through an integrated ethical and institutional success perspective.

Prior studies on AI success and information system performance have predominantly adopted the Information System Success Model (ISSM) and the Technology Acceptance Model (TAM) to explain system usage, satisfaction, and perceived benefits (Al-Hawamleh, 2024). While these models provide strong explanatory power for understanding technology adoption and performance, they were originally designed for organizational and commercial environments, where ethical obligations, public accountability, and citizen trust are not as structurally binding as in government settings (Schmager et al., 2024; Yang et al., 2025). Consequently, these classical models remain insufficient to capture the distinctive ethical, legal, and governance dimensions that inherently characterize public-sector AI implementation.

At the same time, the rapid expansion of Ethical AI frameworks emphasizing principles such as fairness, transparency, accountability, privacy, and human oversight has reshaped how AI governance is conceptualized across public institutions (Tabaghdehi & Ayaz, 2025). However, despite their growing normative importance, these ethical principles are rarely integrated into formal quantitative models of information system success. As a result, the ethical dimension of AI remains largely disconnected from empirical success evaluation, leaving a critical theoretical and methodological gap between AI governance discourse and IS success measurement practices in government contexts (Birkstedt et al., 2023; Krijger, 2022).

Unlike private-sector organizations that primarily evaluate AI success through productivity, profitability, customer satisfaction, and competitive advantage, government institutions operate under broader public accountability obligations. Failures of AI systems in government may not only reduce operational efficiency but can also undermine institutional legitimacy, public trust, transparency, fairness, and democratic accountability. Consequently, evaluating AI success in government requires a framework that extends beyond technical performance and behavioral acceptance by incorporating ethical and governance dimensions. This distinction highlights the need for a dedicated public-sector AI success framework capable of capturing both operational effectiveness and ethical legitimacy.

This study addresses that gap by developing an Ethical Information System Success Framework for Artificial Intelligence in Government, which explicitly integrates ethical AI governance constructs into the causal structure of the IS Success Model (Camilleri, 2024; Xue & Pang, 2022). By embedding ethical fairness, transparency, accountability, and privacy–security into the success mechanism of AI systems, this research reframes AI success as a socio-technical and institutional phenomenon, rather than a purely technical or behavioral outcome. The proposed framework conceptualizes ethical governance as a core antecedent of user trust, satisfaction, actual use, and net public benefits.

Furthermore, the validity and reliability of measurement instruments are central to ensuring the scientific integrity of empirical research. In government AI studies, where ethical legitimacy and policy implications are highly sensitive, instrument rigor becomes even more critical (Grimmelikhuijsen & Meijer, 2022; Martin & Waldman, 2023). Therefore, beyond model



development, this study also focuses on instrument decomposition and triangulation-based validation, ensuring that the ethical–technical constructs are measured with sufficient psychometric and interpretative robustness.

Accordingly, this study is guided by three primary research questions:

- a) How can an ethical-based information system success framework for AI in government be systematically developed?
- b) How can the proposed framework be operationalized into valid and reliable measurement instruments?
- c) To what extent do the ethical and technical constructs empirically explain trust, actual use, and net public benefits of AI in government institutions?

By answering these questions, this study aims to contribute to three main dimensions. Theoretically, it extends the IS Success Model by embedding ethical AI governance as a formal success determinant. Methodologically, it provides a structured procedure for model development, operationalization, and triangulation validation using PLS-SEM and expert judgment. Practically, it offers an evidence-based framework to guide policymakers and government institutions in designing, evaluating, and governing AI systems that are not only efficient, but also ethically legitimate, trustworthy, and publicly accountable.

2. LITERATURE REVIEW AND HYPOTHESIS DEVELOPMENT

Artificial Intelligence (AI) has increasingly become a strategic component of digital government transformation. Governments utilize AI technologies to support public service delivery, policy analysis, administrative automation, predictive analytics, and decision-making processes. The growing dependence on AI has generated substantial interest in understanding the factors that determine successful AI implementation within public-sector organizations. Unlike private-sector environments, government institutions operate under strict requirements for accountability, transparency, legality, and public trust. Consequently, evaluating AI success in government requires a broader perspective that incorporates both technical performance and ethical governance dimensions.

The Information System Success Model (ISSM), originally developed by DeLone and McLean, remains one of the most influential frameworks for evaluating information system effectiveness. The model proposes that system quality, information quality, and service quality influence user satisfaction and system use, which subsequently generate organizational and individual benefits. Numerous studies have confirmed the robustness of ISSM across various technological contexts. However, the model was primarily designed for organizational information systems and does not explicitly address ethical concerns that have become increasingly important in AI-driven environments.

Recent developments in Ethical AI governance emphasize the importance of fairness, transparency, accountability, privacy, and security as fundamental principles for responsible AI deployment. These principles are particularly critical in government contexts because AI systems often influence public services, citizen rights, and policy decisions. Failure to uphold ethical standards may undermine institutional legitimacy and public trust even when technical system performance remains satisfactory. Therefore, integrating Ethical AI principles into the Information System Success framework provides a more comprehensive approach to evaluating AI success in government.

The proposed framework adopts an Input–Process–Output perspective. System Quality, Information Quality, Service Quality, and Ethical AI principles constitute the input factors. User Trust and User Satisfaction represent process variables that shape Actual Use, while Ethical Impact and Net Benefits serve as output variables reflecting the outcomes of AI implementation.

System Quality refers to the degree to which AI systems operate reliably, securely, efficiently, and consistently. Reliable and well-performing systems reduce uncertainty and enhance users'



confidence in technology. Previous studies have demonstrated that system quality positively influences trust because users are more likely to rely on systems that function accurately and consistently. Accordingly, the following hypothesis is proposed:

H1: System Quality positively influences User Trust.

System Quality also contributes to positive user experiences. Systems that are responsive, stable, and easy to use tend to increase user satisfaction because they facilitate task completion and improve work efficiency. Therefore:

H2: System Quality positively influences User Satisfaction.

Information Quality refers to the accuracy, completeness, relevance, timeliness, and consistency of information generated by AI systems. High-quality information enables users to make informed decisions and reduces uncertainty regarding AI outputs. Previous Information System Success studies consistently identify information quality as a critical determinant of trust. Consequently:

H3: Information Quality positively influences User Trust.

Users are more satisfied when information generated by AI is accurate, relevant, and useful for their work activities. Information quality directly affects users' perceptions of system usefulness and value. Thus:

H4: Information Quality positively influences User Satisfaction.

Service Quality represents the quality of support services surrounding AI implementation, including responsiveness, technical assistance, communication, and problem resolution. Effective support services improve users' experiences and reduce frustration when encountering technical issues. Therefore:

H5: Service Quality positively influences User Satisfaction.

Ethical Fairness refers to the extent to which AI systems provide unbiased and non-discriminatory outcomes. Fairness is widely recognized as a core principle of responsible AI because it promotes equal treatment and reduces social inequities. Users are more likely to trust AI systems when they perceive decision-making processes as fair and impartial. Hence:

H6: Ethical Fairness positively influences User Trust.

Ethical Transparency reflects the degree to which AI processes, decision mechanisms, and outputs can be understood, interpreted, and explained. Transparency reduces uncertainty and enables users to evaluate the reasoning behind AI recommendations. Explainable and traceable AI systems are generally perceived as more trustworthy. Therefore:

H7: Ethical Transparency positively influences User Trust.

Ethical Accountability concerns the existence of clear responsibilities, governance mechanisms, auditability, and legal oversight related to AI decisions. Accountability ensures that AI outcomes can be reviewed and that responsible parties can be identified when errors occur. Such mechanisms strengthen institutional legitimacy and increase trust among users. Thus:

H8: Ethical Accountability positively influences User Trust.

Ethical Privacy and Security refer to the protection of personal information, confidentiality, access control, cybersecurity, and responsible data management practices. Government employees



frequently interact with sensitive data, making privacy and security critical determinants of trust. Accordingly:

H9: Ethical Privacy and Security positively influence User Trust.

Trust represents users' confidence in the reliability, security, and ethical behavior of AI systems. Trust plays a central role in reducing uncertainty and encouraging technology usage. Previous research consistently demonstrates that trusted technologies are more likely to be adopted and utilized. Therefore:

H10: User Trust positively influences Actual Use.

User Satisfaction reflects users' overall evaluation of their experiences with AI systems. Satisfied users are more likely to continue using technologies and integrate them into routine work processes. Consequently:

H11: User Satisfaction positively influences Actual Use.

Actual Use represents the extent to which AI systems are regularly and intensively utilized in government operations. According to the Information System Success Model, benefits can only be realized when systems are actively used. Increased utilization enables organizations to capture operational, strategic, and societal value from AI implementation. Therefore:

H12: Actual Use positively influences Ethical Impact and Net Benefits.

Based on the preceding theoretical arguments, the proposed Ethical Information System Success Framework extends the traditional Information System Success Model by incorporating Ethical AI governance dimensions as antecedents of trust. The framework posits that both technical quality factors and ethical governance mechanisms jointly determine user trust and satisfaction, which subsequently drive actual AI use and generate ethical and organizational benefits within government institutions. This integrated perspective recognizes that successful AI implementation in government depends not only on technological performance but also on ethical legitimacy and institutional accountability.

3. METHODS

This study was designed as a model development and validation research employing a sequential mixed-method approach with triangulation-based instrument testing. The methodological framework was constructed to ensure conceptual rigor, empirical robustness, and interpretative validity in developing the Ethical Information System Success Framework for Artificial Intelligence in Government (Braga et al., 2025; Chowdhury & Oredo, 2023). The overall research procedure followed a structured sequence that integrates theoretical synthesis, empirical measurement, and expert-based confirmation to capture both the technical and ethical dimensions of AI success in the public sector.

The AI applications used by respondents included document summarization systems, chatbot-based public service assistants, predictive analytics platforms, intelligent decision-support systems, automated reporting tools, and generative AI applications utilized in administrative processes. These systems were employed across various government functions, including citizen service delivery, internal administration, policy analysis, and data management.

Before detailing the stages of model testing and data analysis, the overall procedural architecture of the research is first illustrated to clarify the logical flow from concept development to empirical validation. The research procedure is presented in Figure 1, which depicts the interconnected stages of preliminary study, model development, operationalization, instrument testing, and final reporting. This procedural overview provides a systematic representation of how the proposed framework was developed and empirically validated.



As illustrated in Figure 1, the study begins with a preliminary theoretical investigation into Artificial Intelligence in government, Ethical AI governance, and the Information System Success Model. The model development phase is conceptually structured through assumption formulation using the Input–Process–Output logic, followed by model adoption from established IS success and ethical AI frameworks, model combination through causal integration, and final adaptation to the government context. This is followed by the operationalization phase in which abstract constructs are translated into measurable indicators and questionnaire items. The instrument testing phase applies a sequential triangulation strategy combining quantitative PLS-SEM analysis and qualitative expert judgment, which is then finalized through confirmation interpretation (Kurtaliqui et al., 2024; Mukhtar et al., 2022) . The last phase consolidates the validated framework into a complete scientific report.

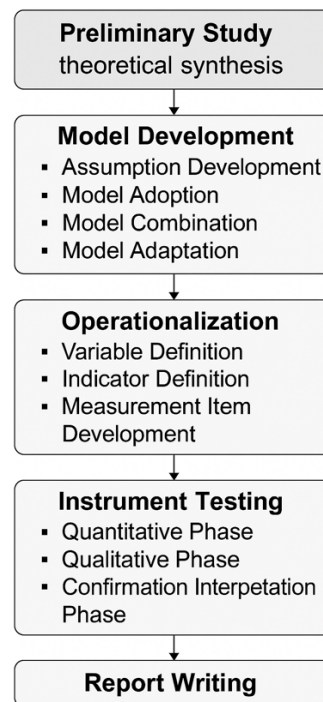


Figure 1 Research Procedure

The model development process yields not only a structural research model but also a conceptual interpretation mechanism that links statistical findings with theoretical justification. This interpretation mechanism is illustrated in Figure 2, which presents the confirmation logic used to integrate quantitative and qualitative validation results. The mechanism provides a systematic basis for evaluating the proposed framework from both empirical and theoretical perspectives.

As illustrated in Figure 2, the statistical results of the measurement and structural models are first interpreted independently. In parallel, expert judgments derived from qualitative interviews are analyzed thematically. These two streams of evidence are then converged through a confirmation interpretation process, in which the empirical findings are cross-validated against established theories of IS success and ethical AI governance. This process ensures that the final research instrument and structural relationships are not only statistically acceptable but also conceptually legitimate within the governance context.

Sample adequacy was assessed using the minimum sample size requirement for PLS-SEM. Following the ten-times rule and inverse square root method, the minimum required sample ranged between 52 and 77 observations, depending on the maximum number of structural paths directed toward an endogenous construct. With 86 valid responses, the sample size exceeded



the minimum requirement and was considered sufficient for model estimation and hypothesis testing.

Following the research design, the operational context of this study involved government employees who actively interact with AI-based systems in their daily administrative and public service tasks. The respondents represent personnel working in digital public services, data analytics units, smart governance operations, and AI-supported decision-making systems. A total of 86 valid responses were obtained using a purposive sampling technique, requiring that participants actively use AI applications in government services, possess at least one year of professional experience in public-sector digital systems, and be directly involved in operational, supervisory, or managerial roles related to AI utilization. Data collection was conducted during the October–November period using an online questionnaire platform, with each response automatically recorded using a timestamp to ensure chronological integrity of the dataset.

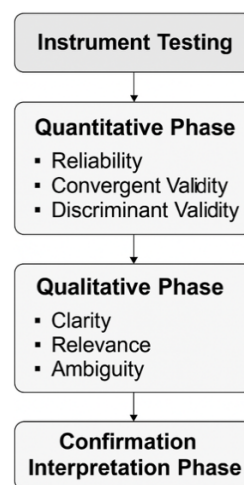


Figure 2 Procedure of the Interpretation Meta-Inference

The research instrument consisted of two integrated components. The first captured the profile of respondents, including age, gender, education level, position, years of service, and institutional type. The second measured the constructs of the proposed ethical-based IS success framework using 55 Likert-scale items distributed across 11 latent variables, namely system quality, information quality, service quality, ethical fairness, ethical transparency, ethical accountability, ethical privacy and security, user trust, user satisfaction, actual use, and net benefits. All measurement items employed a five-point Likert scale, ranging from strongly disagree to strongly agree, to ensure sufficient variance for structural equation modeling.

Quantitative data analysis was conducted using Partial Least Squares Structural Equation Modeling (PLS-SEM) with SmartPLS 4 software, which is particularly suitable for predictive modeling, complex causal structures, and moderate sample sizes. The evaluation of the measurement model focused on indicator reliability through outer loadings, internal consistency reliability through composite reliability values, convergent validity through average variance extracted, and discriminant validity using the Fornell–Larcker criterion. The structural model was assessed through path coefficients, t-values and p-values obtained via bootstrapping procedures, as well as coefficients of determination to evaluate explanatory power and predictive relevance. The significance of structural relationships was evaluated using a bootstrapping procedure with 5,000 subsamples and bias-corrected confidence intervals. This approach enabled robust estimation of t-values, p-values, and confidence intervals for hypothesis testing.

To complement the statistical testing, qualitative validation was conducted through expert judgment involving five specialists in information systems, public administration, AI governance, cybersecurity, and digital policy. Semi-structured interviews were used to assess the clarity of



indicator wording, conceptual relevance of constructs, ethical sufficiency of the framework, and potential sources of bias or ambiguity in the questionnaire. The qualitative findings were then integrated with the quantitative results during the confirmation interpretation stage to refine the final measurement instrument and validate the structural logic of the model.

All research procedures adhered to fundamental ethical principles of social science and information systems research. Participation in the survey was entirely voluntary, informed consent was obtained implicitly through participation, and anonymity was guaranteed for all respondents. No personally identifiable information was disclosed at any stage of the research. Given the sensitive nature of Artificial Intelligence governance in public institutions, particular attention was devoted to data protection, confidentiality, and responsible use of empirical findings.

4. RESULTS AND DISCUSSION

Before conducting the measurement and structural model evaluation, it is important to understand the demographic and professional characteristics of the respondents participating in this study. Respondent profiling provides contextual information regarding the individuals who actively use Artificial Intelligence (AI)-based systems within government institutions and helps establish the representativeness of the collected data. The profile information includes gender, educational background, organizational position, years of professional experience, and institutional affiliation. As presented in Table 1, the respondents represent a diverse range of government employees who are directly involved in the utilization of AI-supported systems in their daily operational and administrative activities.

The respondent profile presented in Table 1 demonstrates that the participants represent a diverse range of government employees actively involved in AI-supported public service activities. Male respondents accounted for 57.0% of the sample, while female respondents represented 43.0%. Most respondents were between 31 and 40 years old and possessed at least a bachelor's degree, indicating a relatively mature and educated workforce. In terms of organizational roles, respondents were distributed across staff, supervisory, and managerial positions, enabling the study to capture perspectives from multiple administrative levels. Furthermore, the majority of respondents had more than six years of professional experience and had been utilizing AI-based systems for at least one year. This profile suggests that the respondents possessed sufficient exposure to government information systems and AI technologies to provide informed evaluations of the proposed framework. Consequently, the respondent composition supports the credibility of the empirical findings and reflects practical experiences associated with AI implementation in government institutions.

Referring to the research procedure, this section presents the results of model development, operationalization, and instrument evaluation for the proposed Ethical Information System Success Framework for Artificial Intelligence in Government. The results are organized sequentially starting from the conceptual assumption of model development, followed by the proposed research model and hypotheses, operational definitions of variables and indicators, measurement instrument construction, and finally the results of statistical and triangulation testing using data from 86 government-sector respondents. Before presenting the operational and statistical results, the theoretical positioning of the proposed framework is first illustrated through the Model Development Assumption and the Developed Research Model and Its Hypotheses, as shown in Figure 3 and Figure 4, respectively. The model assumption illustrates that AI success in government is not only determined by technical qualities of information systems, but also by ethical dimensions that influence trust, satisfaction, and actual use. This reinforces the view that public-sector AI must adhere to transparency, accountability, fairness, and privacy protection to achieve sustainable institutional benefits.

Figure 4 presents the developed research model and its associated hypotheses, which integrate ethical governance explicitly into the structure of the Information System Success model. In this framework, ethical principles are positioned as foundational antecedents that shape institutional trust in government AI systems. Ethical fairness, transparency, accountability, and privacy–



security function as governance mechanisms that influence how users perceive the legitimacy, reliability, and responsibility of AI-based public services. These ethical dimensions, together with core system-related qualities such as system quality, information quality, and service quality, are hypothesized to determine user trust and satisfaction. Trust and satisfaction subsequently influence actual system use, which ultimately leads to the creation of net public benefits. Through this structure, the model conceptualizes ethical governance not merely as a contextual or normative element, but as an active mechanism that generates institutional trust and enables sustainable public value creation through AI utilization.

Table 1 Respondent Profile (n=86)

Category	Classification	Frequency (n)	Percentage (%)
Gender	Male	49	57.0
	Female	37	43.0
Age	21–30 years	18	20.9
	31–40 years	34	39.5
	41–50 years	22	25.6
	> 50 years	12	14.0
Education Level	Diploma	5	5.8
	Bachelor's Degree	41	47.7
	Master's Degree	34	39.5
	Doctoral Degree	6	7.0
Position	Staff	35	40.7
	Supervisor	29	33.7
	Manager/Head of Unit	22	25.6
Years of Service	1–5 years	17	19.8
	6–10 years	25	29.1
	11–15 years	21	24.4
	> 15 years	23	26.7
AI Utilization Experience	Less than 1 year	15	17.4
	1–3 years	38	44.2
	4–5 years	21	24.4
	More than 5 years	12	14.0

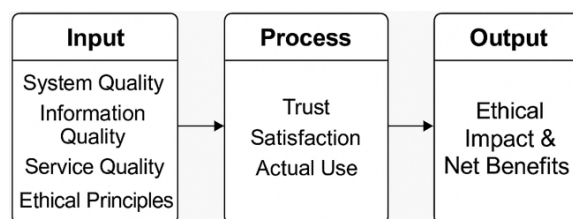


Figure 3 Model Development Assumption (IPO + Ethical AI + IS Success Perspective)

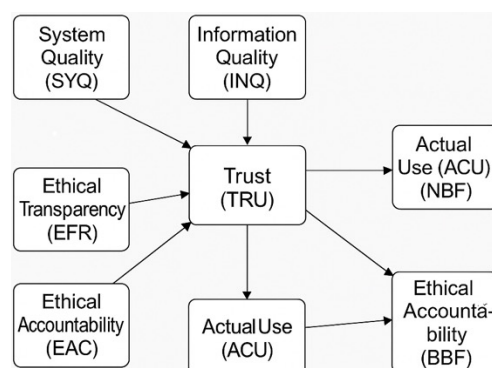


Figure 4 The Developed Research Model and Its Hypotheses



Table 2 presents the definitions of all variables used in this study. Each construct is defined based on established literature in information systems, digital government, and ethical AI governance to ensure conceptual clarity and theoretical consistency. The definitions serve as the foundation for operationalizing variables into measurable indicators, supporting the development of a valid and reliable measurement model for empirical testing. The formal definition of variables ensures that each construct is theoretically grounded in both Information System Success theory and Ethical AI governance literature. Building on these conceptual definitions, the next step involves operationalizing each variable into measurable components that can be empirically assessed.

Table 2 List fo Variable Definitions

Code	Variable	Definition
SYQ	System Quality	The degree to which AI systems operate reliably, securely, and efficiently in government services.
INQ	Information Quality	The degree to which information generated by AI is accurate, timely, and relevant for public decision-making.
SVQ	Service Quality	The degree to which AI service support meets users' expectations in government operations.
EFR	Ethical Fairness	The extent to which AI decisions are free from bias and discrimination.
ETR	Ethical Transparency	The degree to which AI processes and outcomes are explainable and traceable.
EAC	Ethical Accountability	The degree to which responsibility and auditability of AI decisions are clearly defined.
EPS	Ethical Privacy & Security	The degree to which AI protects personal data and sensitive information.
TRU	User Trust	The degree of confidence users place in AI systems in government services.
USF	User Satisfaction	The degree to which users feel satisfied with AI use in their tasks.
ACU	Actual Use	The degree to which AI is routinely used in daily government operations.
NBF	Net Benefits	The overall institutional, ethical, and service performance benefits of AI use.

Accordingly, Table 15 in Appendix A presents the indicators, their conceptual definitions, and the corresponding measurement items used to capture each construct in this study. The indicators were developed through an extensive review of prior IS and AI governance studies and were adapted to the context of artificial intelligence implementation in government. Each measurement item was formulated to reflect observable perceptions and experiences of users when interacting with AI-based public systems, enabling systematic evaluation through a Likert-scale questionnaire. The indicator structure reflects both technical performance and ethical governance dimensions, thereby differentiating this model from conventional Information System Success frameworks.

Because all data were collected using a self-administered questionnaire, common method bias (CMB) was assessed to ensure that the observed relationships were not artificially inflated by the data collection method. Following recommendations in PLS-SEM literature, full collinearity variance inflation factor (VIF) values were examined, with the results presented in Table 3. As shown in Table 3, all full collinearity VIF values are below the recommended threshold of 3.30. These results indicate that common method bias is unlikely to threaten the validity of the findings and that the observed relationships can be interpreted with confidence.

To evaluate the quality of the measurement model, an indicator reliability assessment was conducted using Partial Least Squares–Structural Equation Modeling. This assessment examines the extent to which each indicator consistently represents its underlying construct. The results demonstrate that all constructs exceed the recommended thresholds of composite reliability (CR



≥ 0.70) and average variance extracted ($AVE \geq 0.50$), indicating strong internal consistency reliability and satisfactory convergent validity. Table 4 presents the detailed results of the indicator reliability assessment for all constructs based on 86 valid responses, including the loading values, reliability measures, and acceptance status for each indicator.

Table 3 Common Method Bias Assessment

Construct	Full VIF
SYQ	2.14
INQ	2.08
SVQ	1.95
EFR	1.87
ETR	2.31
EAC	2.28
EPS	2.22
TRU	2.76
USF	2.45
ACU	2.34
NBF	2.11

All measurement indicators demonstrate satisfactory reliability, with outer loading values exceeding the minimum threshold of 0.70. Furthermore, all constructs achieve composite reliability (CR) values above 0.90 and average variance extracted (AVE) values above 0.65, indicating strong internal consistency and adequate convergent validity. Consequently, all indicators are classified as accepted and retained for further structural model analysis. To ensure that each construct is empirically distinct from the others, discriminant validity was subsequently assessed using the Fornell–Larcker criterion. As presented in Table 5, the square root of the AVE for each construct exceeds its correlations with other constructs, confirming that the measurement model satisfies discriminant validity requirements and that each latent variable captures a unique conceptual domain.

Although the Fornell–Larcker criterion remains widely used, recent PLS-SEM guidelines recommend the Heterotrait–Monotrait Ratio (HTMT) as an additional assessment of discriminant validity. Accordingly, HTMT values were calculated and are presented in Table 6. All HTMT values are below the recommended threshold of 0.90, confirming adequate discriminant validity. These findings further support the distinctiveness of the proposed constructs and complement the Fornell–Larcker results.

To strengthen the interpretative validity of the proposed framework, qualitative validation was conducted through expert judgment. The selected experts possessed extensive experience in information systems, public administration, AI governance, cybersecurity, and digital policy, while their profiles and validation focus areas are summarized in Table 7. The expert panel confirmed the conceptual relevance of all major constructs and provided recommendations for improving wording clarity and contextual suitability. The qualitative validation process served as an important complement to statistical testing by ensuring conceptual legitimacy and practical applicability.

Beyond statistical validation, this study applies a triangulation approach to strengthen the robustness of the findings by integrating quantitative results, qualitative expert judgment, and theoretical support. Table 8 presents the triangulation confirmation results, which compare empirical statistical evidence with insights obtained from expert evaluations and established theoretical foundations. This triangulation process ensures that the validated constructs and relationships are not only statistically sound but also conceptually meaningful and practically relevant within the context of ethical AI implementation in government. Expert evaluations consistently identified ethical transparency, accountability, and privacy and security as the most critical dimensions for AI governance in government institutions. Several experts noted that fairness mechanisms are often embedded within technical processes and may not be directly



observable by end users, providing additional explanation for the weaker statistical effect of Ethical Fairness.

Table 4 Results of the Indicator Reliability Assessment with Status (n = 86)

Variable	CR	AVE	Indicator	Outer Loading	Status
System Quality (SYQ)	0.91	0.67	SYQ1	0.81	Accepted
			SYQ2	0.84	Accepted
			SYQ3	0.79	Accepted
			SYQ4	0.83	Accepted
			SYQ5	0.80	Accepted
Information Quality (INQ)	0.92	0.69	INQ1	0.85	Accepted
			INQ2	0.83	Accepted
			INQ3	0.82	Accepted
			INQ4	0.86	Accepted
			INQ5	0.81	Accepted
Service Quality (SVQ)	0.90	0.65	SVQ1	0.79	Accepted
			SVQ2	0.84	Accepted
			SVQ3	0.80	Accepted
			SVQ4	0.78	Accepted
			SVQ5	0.82	Accepted
Ethical Fairness (EFR)	0.93	0.71	EFR1	0.85	Accepted
			EFR2	0.82	Accepted
			EFR3	0.88	Accepted
			EFR4	0.84	Accepted
			EFR5	0.86	Accepted
Ethical Transparency (ETR)	0.95	0.74	ETR1	0.87	Accepted
			ETR2	0.89	Accepted
			ETR3	0.85	Accepted
			ETR4	0.88	Accepted
			ETR5	0.84	Accepted
Ethical Accountability (EAC)	0.94	0.72	EAC1	0.85	Accepted
			EAC2	0.88	Accepted
			EAC3	0.86	Accepted
			EAC4	0.84	Accepted
			EAC5	0.87	Accepted
Ethical Privacy & Security (EPS)	0.92	0.68	EPS1	0.84	Accepted
			EPS2	0.86	Accepted
			EPS3	0.80	Accepted
			EPS4	0.83	Accepted
			EPS5	0.81	Accepted
User Trust (TRU)	0.93	0.70	TRU1	0.86	Accepted
			TRU2	0.83	Accepted
			TRU3	0.85	Accepted
			TRU4	0.87	Accepted
			TRU5	0.82	Accepted
User Satisfaction (USF)	0.91	0.67	USF1	0.84	Accepted
			USF2	0.81	Accepted
			USF3	0.86	Accepted
			USF4	0.80	Accepted
			USF5	0.83	Accepted
Actual Use (ACU)	0.90	0.65	ACU1	0.81	Accepted
			ACU2	0.85	Accepted
			ACU3	0.79	Accepted
			ACU4	0.84	Accepted
			ACU5	0.80	Accepted
Net Benefits (NBF)	0.94	0.73	NBF1	0.87	Accepted
			NBF2	0.85	Accepted
			NBF3	0.88	Accepted
			NBF4	0.84	Accepted
			NBF5	0.86	Accepted



The triangulation process confirms that ethical governance variables—ethical transparency (ETR), ethical accountability (EAC), and ethical privacy and security (EPS)—demonstrate the strongest empirical and conceptual support in the context of government AI implementation. The results provide robust empirical evidence that ethical governance dimensions are not peripheral factors but core determinants of AI success in public institutions. Ethical transparency, accountability, and privacy protection exhibit stronger structural effects on trust and actual system use compared to conventional system, information, and service quality variables. These findings indicate that public-sector AI success is fundamentally shaped by institutional legitimacy and citizen trust rather than solely by technical efficiency or system performance. Table 9 presents the hypothesis testing results, including path coefficients, t-values, p-values, and explanatory power based on data from 86 respondents, which collectively demonstrate the strength and significance of the proposed structural relationships.

Table 5 Discriminant Validity (AVE Square Root Matrix)

Variable	SYQ	INQ	SVQ	EFR	ETR	EAC	EPS	TRU	USF	ACU	NBF
SYQ	0.82										
INQ	0.54	0.83									
SVQ	0.49	0.57	0.81								
EFR	0.41	0.45	0.44	0.84							
ETR	0.43	0.48	0.46	0.63	0.86						
EAC	0.40	0.47	0.45	0.60	0.68	0.85					
EPS	0.44	0.51	0.50	0.58	0.60	0.64	0.83				

Table 6 HTMT Ratio Assessment

Construct Pair	HTMT
SYQ-INQ	0.65
SYQ-TRU	0.52
INQ-TRU	0.58
ETR-TRU	0.74
EAC-TRU	0.71
EPS-TRU	0.68
TRU-ACU	0.81
USF-ACU	0.78

Table 7 Expert Profile and Validation Criteria

Expert	Expertise	Experience
E1	Information Systems	15 years
E2	Public Administration	18 years
E3	AI Governance	12 years
E4	Cybersecurity	14 years
E5	Digital Policy	16 years

Table 8 Triangulation Confirmation Results (Quantitative–Qualitative–Theoretical)

Indicator	Quantitative	Qualitative	Theoretical	Decision
ETR	Accepted	Strong	Strong	Accepted
EAC	Accepted	Strong	Strong	Accepted
EPS	Accepted	Strong	Strong	Accepted
EFR	Revised	Moderate	Strong	Revised
SVQ	Revised	Moderate	Moderate	Revised
INQ	Accepted	Moderate	Strong	Accepted

Beyond direct effects, mediation analysis was conducted to examine whether User Trust functions as an explanatory mechanism linking ethical governance dimensions and actual AI use. The indirect effects are reported in Table 10. The mediation results demonstrate that User Trust



partially mediates the effects of ethical transparency, accountability, and privacy and security on actual AI use. These findings reinforce the argument that ethical governance creates value primarily by strengthening trust toward AI systems.

Table 9 Hypothesis Testing Results (n = 86)

Hypothesis	Path	Std. Beta	t-value	p-value	Decision
H1	SYQ → TRU	0.18	2.41	0.016	Accepted
H2	INQ → TRU	0.21	2.79	0.006	Accepted
H3	EFR → TRU	0.12	1.85	0.065	Rejected (not sig.)
H4	ETR → TRU	0.29	4.03	<0.001	Accepted
H5	EAC → TRU	0.24	3.32	0.001	Accepted
H6	EPS → TRU	0.17	2.27	0.024	Accepted
H7	SYQ → USF	0.23	3.01	0.003	Accepted
H8	INQ → USF	0.27	3.66	<0.001	Accepted
H9	SVQ → USF	0.19	2.58	0.011	Accepted
H10	TRU → ACU	0.39	5.12	<0.001	Accepted
H11	USF → ACU	0.37	4.78	<0.001	Accepted
H12	ACU → NBF	0.80	13.20	<0.001	Accepted

Table 10 Mediation Analysis Results

Indirect Path	Beta	t-value	p-value
ETR→TRU→ACU	0.113	3.61	<0.001
EAC→TRU→ACU	0.094	3.02	0.002
EPS→TRU→ACU	0.066	2.43	0.015

To evaluate the explanatory power of the structural model, the coefficient of determination (R^2) was examined for each endogenous construct. The R^2 values indicate the proportion of variance in the dependent variables that can be explained by their respective predictor constructs. As presented in Table 11, the model demonstrates substantial explanatory power for key endogenous variables, particularly trust, user satisfaction, actual use, and net benefits. High R^2 values suggest that the integrated ethical and system-related predictors effectively account for variations in these outcomes. This finding reinforces the robustness of the proposed framework and confirms that combining ethical governance dimensions with the Information System Success structure provides strong predictive capability for understanding AI implementation outcomes in government settings.

Table 11 Coefficient of Determination (R^2) for Endogenous Constructs

Endogenous Variable	Predictor(s)	R^2	Interpretation
User Trust (TRU)	SYQ, INQ, EFR, ETR, EAC, EPS	0.68	Substantial
User Satisfaction (USF)	SYQ, INQ, SVQ	0.62	Substantial
Actual Use (ACU)	TRU, USF	0.71	Substantial
Net Benefits (NBF)	ACU	0.64	Substantial

To assess the relative contribution of each predictor construct, effect size (f^2) values were calculated, and the results are presented in Table 12. The strongest effect size is observed for the relationship between Actual Use and Net Benefits, confirming that organizational value can only be realized when AI systems are actively utilized. Trust also demonstrates a substantial contribution to actual AI use. Predictive relevance was subsequently evaluated using the blindfolding procedure, with the resulting Q^2 values presented in Table 13. All Q^2 values exceed zero, indicating satisfactory predictive relevance and demonstrating that the model possesses adequate predictive capability.

Finally, overall model fit was assessed using the Standardized Root Mean Square Residual (SRMR), with the result presented in Table 14. The SRMR value of 0.072 falls below the



recommended threshold of 0.08, indicating an acceptable level of model fit. This result provides additional evidence supporting the adequacy of the proposed Ethical Information System Success Framework.

Table 12 Effect Size (f^2)

Path	f^2	Interpretation
ETR→TRU	0.22	Medium
EAC→TRU	0.16	Medium
EPS→TRU	0.10	Small
TRU→ACU	0.34	Large
USF→ACU	0.28	Medium
ACU→NBF	0.89	Large

Table 13 Predictive Relevant (Q^2)

Construct	Q^2
TRU	0.47
USF	0.42
ACU	0.51
NBF	0.56

Table 14 Model Fit Assessment

Index	Value	Threshold
SRMR	0.072	< 0.08

The structural model results indicate that most hypothesized relationships are statistically significant. System Quality (H1) and Information Quality (H2) both have positive and significant effects on User Trust, confirming that reliable and accurate AI systems are essential foundations for trust formation in government settings. Ethical Transparency (H4), Ethical Accountability (H5), and Ethical Privacy & Security (H6) show relatively strong and significant paths to User Trust, highlighting the central role of ethical governance in building confidence toward AI in government. Ethical Fairness (H3), although positive, does not reach the conventional significance level, suggesting that fairness is perceived as important but may be less salient or less visible to internal government users compared to transparency and accountability mechanisms.

Regarding User Satisfaction, System Quality (H7), Information Quality (H8), and Service Quality (H9) all significantly influence satisfaction levels. This confirms that user satisfaction is shaped by a combination of technical robustness, informational value, and service support surrounding AI use (Al-Besher & Kumar, 2022; Kim & Kim, 2025). The relatively higher coefficient of Information Quality indicates that, in government environments, the perceived usefulness and reliability of AI-generated information are particularly critical for users' satisfaction judgments.

The downstream effects of User Trust and User Satisfaction on Actual Use are also substantial and statistically significant (H10 and H11). Both constructs jointly explain 71% of the variance in Actual Use ($R^2 = 0.71$), indicating that ethical and experiential perceptions are key drivers of AI adoption in routine government operations. Finally, Actual Use has a very strong and significant effect on Net Benefits (H12), with a path coefficient of 0.80 and explaining 64% of its variance. This supports the classical IS Success logic that realized benefits emerge only when systems are actually and intensively used, now extended with an ethical AI perspective in the government sector (Alhosani & Alhashmi, 2024; Schmager et al., 2024).

From a theoretical standpoint, the integration of Ethical AI principles into the IS Success Model successfully extends the DeLone and McLean framework into the public governance context. Traditional IS success research emphasizes usability, satisfaction, and performance, while this study demonstrates that ethical legitimacy acts as a prerequisite for technology acceptance and



sustained use in government organizations. This finding reinforces contemporary views that digital transformation in government must be framed within responsible AI governance (Birkstedt et al., 2023; Pappas et al., 2023).

Methodologically, this study confirms the robustness of combining PLS-SEM statistical validation with qualitative expert judgment through triangulation. With a sample size of 86 respondents, the outer model demonstrates strong psychometric properties. The triangulation results further refine the instrument by filtering statistically weak but theoretically important indicators. This confirms that ethical constructs require both statistical rigor and normative validation (Braga et al., 2025; Camilleri, 2024; Chowdhury & Oredo, 2023; Lambert & Newman, 2023; Xue & Pang, 2022).

One of the most interesting findings is the non-significant effect of Ethical Fairness on User Trust. Although fairness is widely recognized as a fundamental principle of responsible AI, government employees may have limited visibility into algorithmic fairness mechanisms embedded within AI systems. In contrast, transparency and accountability are directly observable through explanations, audit trails, reporting procedures, and governance policies. Consequently, users may evaluate AI trustworthiness primarily through visible governance mechanisms rather than through perceptions of algorithmic fairness. This finding suggests that fairness remains an important normative principle but may operate indirectly through transparency and accountability structures in government AI environments.

Practically, the findings imply that governments implementing AI should prioritize ethical governance infrastructure alongside technical deployment. Establishing auditability mechanisms, explainable AI interfaces, and privacy-by-design architectures will significantly enhance user trust and institutional benefits. Without these, technical excellence alone is insufficient to generate sustainable net public value from AI systems.

5. CONCLUSIONS

This study developed and validated an Ethical Information System Success Framework for Artificial Intelligence in Government by integrating Ethical AI governance principles with the Information System Success Model. The findings demonstrate that the success of AI implementation in government institutions is influenced not only by traditional information system quality dimensions but also by ethical governance factors. In particular, ethical transparency, accountability, and privacy–security were found to significantly influence user trust, which subsequently contributes to actual AI use and the realization of net public benefits. In contrast, ethical fairness did not exhibit a statistically significant direct effect on user trust, suggesting that fairness may be perceived differently or indirectly through other governance mechanisms within government environments.

The primary theoretical contribution of this study lies in extending the classical Information System Success Model by incorporating ethical governance constructs as formal antecedents of AI success. The findings provide empirical evidence that ethical governance should be viewed not merely as a compliance requirement or normative guideline, but as an integral determinant of trust formation, system utilization, and organizational outcomes. By positioning ethical transparency, accountability, and privacy–security within the success mechanism of AI systems, this study contributes to the growing literature on responsible AI and digital government.

From a practical perspective, the proposed framework offers a structured approach for evaluating and governing AI implementation in public-sector organizations. The findings suggest that government institutions should prioritize explainability, auditability, accountability mechanisms, and privacy protection alongside technical system development. Investments in ethical governance infrastructure are essential for strengthening user trust, encouraging sustained AI utilization, and maximizing institutional and public-service benefits.

Despite these contributions, several limitations should be acknowledged. First, the study utilized a relatively moderate sample size of 86 respondents drawn from government employees in



Indonesia, which may limit the generalizability of the findings to other national and institutional contexts. Second, the research relied on self-reported perceptions rather than objective organizational performance indicators. Third, the cross-sectional design limits the ability to examine changes in trust, satisfaction, and AI utilization over time. Finally, the study focused primarily on internal government users and did not include citizen perspectives, which represent an important stakeholder group in public-sector AI implementation.

Future research is encouraged to validate the proposed framework using larger and more diverse samples across different countries and administrative settings. Subsequent studies may incorporate citizens, policymakers, and external stakeholders to obtain a broader understanding of AI success in government. Longitudinal research designs could provide deeper insights into the evolution of trust and AI adoption over time. In addition, future investigations may integrate objective performance indicators, organizational culture, regulatory maturity, and emerging ethical dimensions to further refine and strengthen the proposed Ethical Information System Success Framework.

REFERENCES

- Al-Ansi, A. M., Garad, A., Jaboob, M., & Al-Ansi, A. (2024). Elevating e-Government: Unleashing the Power of AI and IoT for Enhanced Public Services. *Heliyon*, 10(23), Article ID: e40591. <https://doi.org/10.1016/j.heliyon.2024.e40591>
- Al-Besher, A., & Kumar, K. (2022). Use of Artificial Intelligence to Enhance e-Government Services. *Measurement: Sensors*, 24, Article ID: 100484. <https://doi.org/10.1016/j.measen.2022.100484>
- Al-Dulaimi, A. O. M., & Mohammed, M. A.-A. W. (2026). Legal Responsibility for Errors Caused by Artificial Intelligence (AI) in the Public Sector. *International Journal of Law and Management*, 68(4), 695–722. <https://doi.org/10.1108/IJLMA-08-2024-0295>
- Al-Hawamleh, A. (2024). Exploring the Satisfaction and Continuance Intention to Use E-Learning Systems. *International Journal of Electrical and Computer Engineering Systems*, 15(2), 201–214. <https://doi.org/10.32985/ijeces.15.2.8>
- Alhosani, K., & Alhashmi, S. M. (2024). Opportunities, Challenges, and Benefits of AI Innovation in Government Services: A Review. *Discover Artificial Intelligence*, 4(1), Article ID: 18. <https://doi.org/10.1007/s44163-024-00111-w>
- Anshari, M., Hamdan, M., Ahmad, N., & Ali, E. (2025). Public Service Delivery, Artificial Intelligence and the Sustainable Development Goals: Trends, Evidence and Complexities. *Journal of Science and Technology Policy Management*, 16(1), 163–181. <https://doi.org/10.1108/JSTPM-07-2023-0123>
- Birkstedt, T., Minkkinen, M., Tandon, A., & Mäntymäki, M. (2023). AI Governance: Themes, Knowledge Gaps and Future Agendas. *Internet Research*, 33(7), 133–167. <https://doi.org/10.1108/INTR-01-2022-0042>
- Braga, C. M., Serrano, M. A., & Fernández-Medina, E. (2025). Towards a Methodology for Ethical Artificial Intelligence System Development: A Necessary Trustworthiness Taxonomy. *Expert Systems with Applications*, 286, Article ID: 128034. <https://doi.org/10.1016/j.eswa.2025.128034>
- Camilleri, M. A. (2024). Artificial Intelligence Governance: Ethical Considerations and Implications for Social Responsibility. *Expert Systems*, 41(7), Article ID: e13406. <https://doi.org/10.1111/exsy.13406>
- Chowdhury, T., & Oredo, J. (2023). AI Ethical Biases: Normative and Information Systems Development Conceptual Framework. *Journal of Decision Systems*, 32(3), 617–633. <https://doi.org/10.1080/12460125.2022.2062849>
- Grimmelikhuijsen, S., & Meijer, A. (2022). Legitimacy of Algorithmic Decision-Making: Six Threats and the Need for a Calibrated Institutional Response. *Perspectives on Public Management and Governance*, 5(3), 232–242. <https://doi.org/10.1093/ppmgov/gvac008>
- Tabaghdehi, S. A. H., & Ayaz, Ö. (2025). AI Ethics in Action: A Circular Model for Transparency, Accountability and Inclusivity. *Journal of Managerial Psychology*. <https://doi.org/10.1108/JMP-03-2024-0177>



- Kim, S. Y., & Kim, J. (2025). The Impact of AI Recommendation Quality on Service Satisfaction: The Moderating Roles of Standardization and Customization. *Journal of Services Marketing*, 39(4), 365–386. <https://doi.org/10.1108/JSM-05-2024-0214>
- Krijger, J. (2022). Enter the Metrics: Critical Theory and Organizational Operationalization of AI Ethics. *AI & SOCIETY*, 37(4), 1427–1437. <https://doi.org/10.1007/s00146-021-01256-3>
- Kurtaliqi, F., Lancelot Miltgen, C., Viglia, G., & Pantin-Sohier, G. (2024). Using Advanced Mixed Methods Approaches: Combining PLS-SEM and Qualitative Studies. *Journal of Business Research*, 172, Article ID: 114464. <https://doi.org/10.1016/j.jbusres.2023.114464>
- Lambert, L. S., & Newman, D. A. (2023). Construct Development and Validation in Three Practical Steps: Recommendations for Reviewers, Editors, and Authors. *Organizational Research Methods*, 26(4), 574–607. <https://doi.org/10.1177/10944281221115374>
- Martin, K., & Waldman, A. (2023). Are Algorithmic Decisions Legitimate? The Effect of Process and Outcomes on Perceptions of Legitimacy of AI Decisions. *Journal of Business Ethics*, 183(3), 653–670. <https://doi.org/10.1007/s10551-021-05032-7>
- Mukhtar, N., Kamin, Y. Bin, & Saud, M. S. B. (2022). Quantitative Validation of a Proposed Technical Sustainability Competency Model: A PLS-SEM Approach. *Frontiers in Sustainability*, 3, Article ID: 841643. <https://doi.org/10.3389/frsus.2022.841643>
- Pappas, I. O., Mikalef, P., Dwivedi, Y. K., Jaccheri, L., & Krogstie, J. (2023). Responsible Digital Transformation for a Sustainable Society. *Information Systems Frontiers*, 25(3), 945–953. <https://doi.org/10.1007/s10796-023-10406-5>
- Schmager, S., Grøder, C. H., Parmiggiani, E., Pappas, I., & Vassilakopoulou, P. (2024). Exploring Citizens' Stances on AI in Public Services: A Social Contract Perspective. *Data & Policy*, 6, Article ID: e19. <https://doi.org/10.1017/dap.2024.13>
- Vatamanu, A. F., & Tofan, M. (2025). Integrating Artificial Intelligence into Public Administration: Challenges and Vulnerabilities. *Administrative Sciences*, 15(4), Article ID: 149. <https://doi.org/10.3390/admsci15040149>
- Xue, L., & Pang, Z. (2022). Ethical Governance of Artificial Intelligence: An Integrated Analytical Framework. *Journal of Digital Economy*, 1(1), 44–52. <https://doi.org/10.1016/j.jdec.2022.08.003>
- Yang, Y., Peng, L., & Wan, D. (2025). A Comparative Study on the Acceptance of Autonomous Driving Technology by China and Europe: A Cross-Cultural Empirical Analysis Based on the Technology Acceptance Model. *World Electric Vehicle Journal*, 16(11), Article ID: 589. <https://doi.org/10.3390/wevj16110589>
- Yigitcanlar, T., David, A., Li, W., Fookes, C., Bibri, S. E., & Ye, X. (2024). Unlocking Artificial Intelligence Adoption in Local Governments: Best Practice Lessons from Real-World Implementations. *Smart Cities*, 7(4), 1576–1625. <https://doi.org/10.3390/smartcities7040064>



APPENDIX A

Table 15 Indicators, Definitions, and Measurement Items

Variable	Code	Indicator	Indicator Definition	Measurement Item
System Quality (SYQ)	SYQ1	Reliability	The degree to which the AI system operates consistently without failures.	The AI system used in my institution operates reliably.
	SYQ2	Availability	The extent to which the AI system is accessible whenever needed.	The AI system is always available when required for my work.
	SYQ3	Security	The degree to which the AI system is protected from unauthorized access.	The AI system is secure from unauthorized access and misuse.
	SYQ4	Integration	The extent to which the AI system integrates well with other government systems.	The AI system integrates smoothly with other government information systems.
	SYQ5	Response Speed	The degree to which the AI system responds quickly to requests.	The AI system responds quickly to my requests.
Information (INQ)	INQ1	Accuracy	The extent to which AI generates correct and precise information.	The information provided by AI is accurate for my work.
	INQ2	Timeliness	The degree to which AI provides up-to-date information.	The AI provides updated information when needed.
	INQ3	Relevance	The degree to which AI information is relevant to government tasks.	The information from AI is relevant to my job tasks.
	INQ4	Completeness	The extent to which AI information covers all necessary aspects.	The information from AI is complete for decision-making.
	INQ5	Consistency	The degree to which AI outputs are stable and reliable over time.	The information from AI is consistent across different uses.
Service Quality (SVQ)	SVQ1	Responsiveness	The speed of response of AI-related services to user issues.	AI services respond quickly to my problems.
	SVQ2	Technical Support	The adequacy of technical assistance for AI users.	I receive sufficient technical support when using AI.
	SVQ3	Service Reliability	The dependability of AI service delivery.	AI services are reliable when I need help.
	SVQ4	Communication	The clarity of communication from AI service providers.	Communication related to AI services is clear and understandable.



Variable	Code	Indicator	Indicator Definition	Measurement Item
	SVQ5	Problem Resolution	The effectiveness of AI services in solving user problems.	AI services resolve my problems effectively.
Ethical Fairness (EFR)	EFR1	Non-discrimination	AI does not discriminate against specific groups.	AI treats all users fairly without discrimination.
	EFR2	Algorithmic Neutrality	AI decisions are not politically or socially biased.	AI decisions are neutral and unbiased.
	EFR3	Equal Treatment	AI handles similar cases in similar ways.	AI provides equal treatment for similar cases.
	EFR4	Bias Control	AI bias is monitored and mitigated.	There are mechanisms to control bias in AI.
	EFR5	Fair Outcome	AI generates socially fair outcomes.	AI decisions result in fair outcomes.
Ethical Transparency (ETR)	ETR1	Explainability	AI outputs can be explained logically.	AI decisions can be clearly explained.
	ETR2	Traceability	AI decisions can be traced to their data sources.	AI decisions can be traced back to their data sources.
	ETR3	Interpretability	AI behavior can be interpreted by humans.	I can understand how AI produces its outputs.
	ETR4	Audit Access	AI decisions are open to audit.	AI decision logs can be audited when needed.
	ETR5	Decision Clarity	AI recommendations are clearly presented.	AI outputs are clearly presented for decision-making.
Ethical Accountability (EAC)	EAC1	Responsibility Clarity	Responsibility for AI decisions is clearly defined.	There is clear responsibility for AI decisions.
	EAC2	Legal Compliance	AI use complies with laws and regulations.	AI use complies with applicable laws and regulations.
	EAC3	Auditability	AI processes can be audited properly.	AI processes can be audited systematically.
	EAC4	Error Accountability	AI errors can be clearly accounted for.	There are clear procedures for handling AI errors.
	EAC5	Governance Control	AI use is governed by formal regulations.	AI use is regulated by formal governance policies.
Ethical Privacy & Security (EPS)	EPS1	Data Protection	AI protects personal and sensitive data.	AI protects personal and sensitive data.
	EPS2	Confidentiality	AI prevents unauthorized data disclosure.	AI maintains the confidentiality of sensitive information.
	EPS3	Access Control	AI access is restricted to authorized users.	Only authorized users can access the AI system.
	EPS4	Encryption	AI data is encrypted during processing and storage.	AI data is encrypted for security purposes.
	EPS5	Cybersecurity	AI is protected from cyberattacks.	The AI system is protected from cyberattacks.
User Trust (TRU)	TRU1	Reliability Trust	Trust based on system reliability.	I trust AI because it works reliably.
	TRU2	Security Trust	Trust based on system security.	I trust AI because it is secure.
	TRU3	Ethical Trust	Trust based on ethical behavior of AI.	I trust AI because it behaves ethically.



Variable	Code	Indicator	Indicator Definition	Measurement Item
User Satisfaction (USF)	TRU4	System Dependability	Trust in the dependability of AI.	I can depend on AI to support my work.
	TRU5	Institutional Confidence	Trust in the institution due to AI use.	AI increases my confidence in my institution.
	USF1	Overall Satisfaction	Overall satisfaction with AI use.	Overall, I am satisfied with the AI system.
	USF2	Functional Satisfaction	Satisfaction with AI functions.	I am satisfied with the functions of the AI system.
	USF3	Performance Satisfaction	Satisfaction with AI performance.	I am satisfied with the performance of the AI system.
Actual Use (ACU)	USF4	Support Satisfaction	Satisfaction with AI support services.	I am satisfied with AI technical support.
	USF5	Ethical Satisfaction	Satisfaction with Ethical AI behavior.	I am satisfied with the ethical behavior of AI.
	ACU1	Frequency	How often AI is used.	I use AI frequently in my daily work.
	ACU2	Duration	Length of AI use per session.	I use AI for long periods during my work.
	ACU3	Intensity	Depth of AI use.	I use AI intensively for complex tasks.
Net Benefits (NBF)	ACU4	Task Integration	AI is embedded in workflows.	AI is integrated into my daily work processes.
	ACU5	Service Dependency	Dependency of services on AI.	My work depends heavily on AI systems.
	NBF1	Service Effectiveness	AI improves public service quality.	AI improves the effectiveness of public services.
	NBF2	Public Trust	AI increases citizen trust.	AI increases public trust in government services.
	NBF3	Transparency Improvement	AI improves government transparency.	AI increases transparency in government processes.
	NBF4	Decision Quality	AI improves decision-making quality.	AI improves the quality of government decisions.
	NBF5	Institutional Performance	AI improves organizational performance.	AI improves my institution's overall performance.

