

Ensuring Human-in-the-loop Integrity: A Strategic Hr Framework for Ai-driven Compliance in Federal Agencies

Farzana Afroz^{*}, Tamanna Jahan, Ashim Sen Gupta, Rokhshana Parveen and Shuvashish Roy

¹MS in Human Resources, University of New Haven, Connecticut, USA

²MS in Human Resources, University of New Haven, Connecticut, USA

³First Assistant Vice President, International Division, Social Islami Bank PLC, Bangladesh

⁴Doctor of Business Administration and Management, University of Potomac, Washington, D.C., USA

⁵Senior Researcher, Research & Innovation Division, Prime Bank PLC, Dhaka, Bangladesh

*Correspondence: <afroz6bd@gmail.com>

Published: 2026

DOI: <https://doi.org/10.48165/dbitdj.2026.3.01.01>

ABSTRACT

The rapid adoption of AI in public administration is altering decision-making processes and service delivery across government institutions. Federal agencies increasingly use AI systems for tasks such as fraud detection, benefit administration and public service management. Although these technologies are offering innovative efficiency, they also raise concerns regarding accountability, transparency and regulatory compliance. Several public organizations still face challenges in ensuring effective human oversight and governance of AI systems. So, how organizational governance mechanisms can strengthen AI-driven compliance effectiveness in government institutions have been examined in this work. The main purpose of the research is to analyze the effect of three organizational factors on AI compliance effectiveness: Human-in-the-Loop (HITL) oversight mechanisms, governance of AI infrastructure and strategic human resource capability for AI oversight. A quantitative research design using a structured survey instrument has been applied in the research. Data are analyzed using PLS-SEM to examine the relationships among the variables and to evaluate the validity and reliability of the measurement model. The findings confirm all three independent variables significantly influence AI compliance effectiveness. Human-in-the-Loop oversight has a positive strong effect, which indicates that structured human supervision improves accountability in algorithmic decision-making. AI governance infrastructure also demonstrates a significant connection with compliance effectiveness. Additionally, Strategic human resource capability strengthens compliance outcomes. The findings suggest that successful AI governance requires strong human oversight, clear institutional policies and skilled personnel who can manage algorithmic risks. Managers and policymakers should therefore develop formal AI governance frameworks, strengthen human supervision systems and invest in workforce training on AI ethics and compliance. Despite its contributions, the study has some limitations since cross-sectional design and reliance on survey responses were used.

Keywords: Artificial Intelligence Governance; Human-in-the-Loop Oversight; AI Compliance Effectiveness; Public Sector AI Management; Strategic Human Resource Capability; Algorithmic Accountability; AI Governance Infrastructure; Responsible AI Deployment

INTRODUCTION

Federal agencies increasingly adopt artificial intelligence systems to support administrative and regulatory functions to assist in fraud detection, benefit adjudication, law enforcement analytics and customer service operations(Young et al., 2022). Agencies expect AI tools to improve efficiency and data processing capacity. However, rapid technological adoption also raises serious concerns

about privacy protection, algorithmic bias and regulatory compliance(Olawade et al., 2025). It is warned by policy researchers that poorly governed AI systems may weaken public trust and may reinforce discriminatory patterns in decision-making processes (Ada Lovelace Institute, 2021). Therefore, international governance frameworks emphasize the necessity for human-centered and trustworthy in the public sector's AI systems.

Several international institutions have already established guidelines to ensure responsible AI use. The Organization for Economic Co-operation and Development introduced the **AI Principles**, which stress the importance of human agency and continuous oversight in all stages of AI development and arrangement (OECD, 2023). Similarly, the European Union published its **Trustworthy AI guidelines**, which require that humans remain informed and empowered when interacting with automated systems. However, these guidelines encourage governance models such as human-in-the-loop (HITL) and human-on-the-loop oversight to maintain accountability and transparency in algorithmic decision making (European Union, 2024). In the United States, the **AI Risk Management Framework** has been developed by the **National Institute of Standards and Technology**, which emphasizes trustworthiness, ethical accountability and continuous monitoring throughout the AI lifecycle (NIST, 2021).

Despite these policy commitments, implementation gaps are still here significantly in practice. Several studies indicate that the presence of a human reviewer does not always guarantee effective oversight. Human supervisors frequently rely on algorithmic recommendations when they evaluate automated outputs. This behavior reduces the independence of human judgment and weakens accountability mechanisms (Levinson-Waldman & Reynolds, 2024). A systematic review of public-sector AI governance also identifies transparency, explainability and human oversight as the most critical institutional safeguards for responsible automation (TOLEDO, 2026). At the same time, human resource scholars emphasize that technological transformation within organizations often overlooks the human dimension. Research by Fenwick et al. demonstrates that organizations that ignore worker training and workforce adaptation may undermine trust and slow the implementation of AI-enhanced systems (Fenwick et al., 2024). Therefore, the integration of AI governance and human resource management has become a critical issue in modern public administration.

Recent government data illustrate the rapid growth of AI adoption across federal agencies in the U.S. The U.S. Department of Justice reported **315 documented AI use cases in 2025**, which represents an increase of **30.7 percent compared with 2024** (U.S. Department of Justice, 2025). These systems support functions such as investigative analytics, case management and operational intelligence. A similar expansion appears in the Department of Homeland Security. The department's AI inventory increased from **67 use cases to 158 active applications**, which indicates rapid institutional reliance on algorithmic tools for security and administrative tasks (Levinson-Waldman & Reynolds, 2024). The Internal

Revenue Service also reported **126 active AI projects in 2025**, mainly focused on tax administration and compliance monitoring (U.S. Government Accountability Office, 2026).

However, several policy evaluations reveal important governance weaknesses. The U.S. Government Accountability Office reported that some agencies maintain incomplete AI inventories and lack systematic oversight mechanisms. For example, the GAO found that the IRS did not include benefit descriptions for approximately **25 percent of its AI tools**, which limits the agency's ability to evaluate program effectiveness and compliance risks (U.S. Government Accountability Office, 2026). The review also indicated that the IRS lost a significant portion of its AI-trained workforce, which further reduced the agency's oversight capacity (U.S. Government Accountability Office, 2026). These findings indicate that many agencies have expanded their technological infrastructure without developing adequate institutional mechanisms to supervise automated decision systems.

These developments highlight the growing importance of structured governance frameworks. Federal agencies must ensure that human reviewers possess the technical skills, institutional authority and organizational support necessary to supervise algorithmic systems effectively. Oversight responsibilities require expertise in data analysis, algorithm auditing and ethical risk assessment. At the same time, agencies must ensure that AI investments remain aligned with public service objectives and regulatory requirements (U.S. Government Accountability Office, 2026). Without these institutional capacities, automated decision systems may generate operational risks and legal challenges.

This research addresses these challenges by inspecting the role of strategic HR management in strengthening AI governance within federal agencies. The study integrates perspectives from AI governance, public administration compliance and human resource management to develop a people-centered oversight model. The existing policy frameworks emphasize the importance of human supervision in automated systems (OECD, 2023; European Union, 2024). However, empirical studies and government evaluations demonstrate a important gap between strategy commitments and organizational application (Levinson-Waldman & Reynolds, 2024; U.S. Government Accountability Office, 2026). Therefore, this study proposes a strategic HR framework that can support human-in-the-loop integrity and enhance accountability in federal AI deployments.

Problem Statement

The federal government of the United States has been being rapidly expanded the use of AI and systemsthat automate the decision in public administration to support functions such as law enforcement analytics, tax administration, immigration monitoring and policy evaluation(Szucs, 2026). Federal agencies use AI to increase efficiency and to improve data analysis(Kelly, 2026). However, the expansion of AI systems is being created a governance challenge. Many agencies deploy automated tools, yet human oversight systems are found weak and inconsistent. As a result, concerns are raised about accountability, transparency and compliance with federal AI governance standards(Sony et al., 2025).

Recent data shows the scale of AI adoption across federal agencies. The U.S. Department of Justice maintains a list of official AI inventory consisting 315 AI use cases in 2025 (U.S. Department of Justice, 2025). These applications support tasks such as predictive analytics, case management and investigative intelligence. The department established a AI Officer (Chief) and an Emerging Technology Board to coordinate governance. However, the oversight structure depends largely on voluntary coordination among offices(OECD, 2025b). Limited staffing and limited specialized expertise may reduce the board’s capacity to supervise all AI deployments effectively (U.S. Department of Justice, 2025).

A similar challenge is found in the Internal Revenue Service. The agency reported 126 active AI projects in 2025 that support tax administration and fraud detection including data matching, compliance prediction and automated risk scoring(OECD, 2025a). However, the Officeof U.S. Government Accountability found several governance gaps in the IRS AI program. The GAO reported that the agency’s AI inventory was incomplete and did not include all operational systems. The review also found that the IRS lost approximately 20 percent of its AI-skilled personnel between 2024 and 2025. This decline in specialized staff reduced the agency’s capacity to supervise algorithmic systems and weakened the integration of human oversight processes (U.S. Government Accountability Office, 2026).

The situation is also an evident in the Department of Homeland Security. According to policy reviews, the department reported 158 active AI use cases across its operational components (Levinson-Waldman & Reynolds, 2024). These systems support border security analysis, immigration enforcement and threat detection. Oversight responsibilities are seen within the office of the Chief Technology Officer and within agency privacy boards. However, the department frequently categorizes systems as

“low risk” when a human simply reviews the output of the algorithm. Though several research indicates that this approach does not always ensure meaningful oversight because human reviewers often rely heavily on algorithmic recommendations (Levinson-Waldman & Reynolds, 2024). As a result, the presence of a human reviewer does not always guarantee accountability or independent judgment.

The variation across agencies illustrates a broader governance problem. Federal policies require responsible AI use like Executive Order 13960 established principles for trustworthy AI in government operations and Executive Order 14110 strengthened federal AI governance standards(The White House, 2020). In addition, these policies require transparency, risk management and human oversight in automated systems. However, many agencies have lack of standardized human resource mechanisms to implement these principles effectively. In particular, agencies have lack of structured processes for recruiting AI governance specialists, training employees in algorithmic auditing and organizing interdisciplinary teams.

Table 1: Growing number of AI applications in their operations			
Agency	Active AI Use Cases (Year)	Oversight Structure	Human Oversight Practice
DOJ	315 (2025) (U.S. Department of Justice, 2025)	Chief AI Officer; Emerging Technology Board (U.S. Department of Justice, 2025)	Board reviews and consolidates AI use cases; human managers participate in operational decisions
IRS	126 (2025) (U.S. Government Accountability Office, 2026)	Office of AI Governance	GAO identified incomplete AI inventory and loss of 20% of AI-skilled staff between 2024–2025; oversight structure still developing (U.S. Government Accountability Office, 2026)
DHS	158 (2025) (Levinson-Waldman & Reynolds, 2024)	Chief Technology Officer for AI; ICE privacy boards	Systems often categorized as low risk when human review exists; risk evaluation may remain limited (Levinson-Waldman & Reynolds, 2024)

The table (table 1) shows that federal agencies recognize the growing number of AI applications in their operations. However, their approaches to human oversightisvery different. In many situations, agencies assume that the presence of a human reviewer satisfies governance requirements and that’s why this assumption may create a

false sense of accountability because humans often defer to algorithmic recommendations rather than challenge them (Levinson-Waldman & Reynolds, 2024). Therefore, the human-in-the-loop concept oversight demands more than simple human presence, structured procedures, specialized training and clear accountability mechanisms.

There are some policy reviews which also highlight structural challenges in managing AI investments and oversight systems. The U.S. Government Accountability Office recommended that the IRS improve coordination across divisions and strengthen management practices for AI programs (U.S. Government Accountability Office, 2026). These recommendations indicate that the problem is beyond technology management, organizational capacity, workforce development and governance design.

The problem therefore exists at two interconnected levels. The first level is technical governance where agencies often have lack of standardized metrics for algorithmic evaluation, transparent audit trails and systematic risk assessments for AI systems. Without these mechanisms, it becomes difficult to verify whether automated decisions comply with legal and ethical standards. The second level is organizational capacity where human resource departments in many agencies are not yet prepared to manage the workforce transformation required by AI governance. Agencies must recruit specialists in data science, algorithm auditing and AI ethics. Moreover, they must also train existing employees to understand algorithmic outputs and to exercise independent oversight.

This gap between technological deployment and human oversight creates serious compliance risks. Automated systems can influence decisions that affect legal rights, financial outcomes and public trust. When agencies havenot-sufficient structured human-in-the-loop processes, they may fail to identify algorithmic bias, operational errors, or unintended consequences. Therefore, federal agencies require a strategic human resource framework that integrates AI governance with workforce management. Such a framework must support recruitment of specialized staff, training of oversight teams and institutional mechanisms that ensure meaningful human control over automated decisions.

The absence of a coordinated HR strategy for AI oversight represents a critical governance challenge in the federal government. Addressing this gap is essential for ensuring responsible AI deployment and for maintaining compliance with emerging federal AI regulations. Consequently, a Strategic HR Framework for Human-in-the-Loop Integrity is needed to strengthen AI-driven compliance mechanisms across federal agencies.

Objectives

- ❑ To examine the impact of human-in-the-loop oversight mechanisms on AI-driven compliance effectiveness.
- ❑ To analyze the role of AI governance infrastructure in ensuring accountable AI deployment.
- ❑ To evaluate how strategic HR capability supports oversight and compliance in AI systems.
- ❑ To develop a strategic HR framework that strengthens human-in-the-loop integrity in federal AI governance.

Research Questions

1. How do human-in-the-loop oversight mechanisms influence AI-driven compliance effectiveness in federal agencies?
2. How does AI governance infrastructure contribute to accountable AI deployment?
3. How does strategic HR capability support human oversight and AI compliance?
4. How can strategic HR policies strengthen human-in-the-loop integrity in AI governance?

LITERATURE REVIEW

Research on artificial intelligence in public administration frequently highlights a balance between efficiency gains and accountability requirements. AI technologies are being adopted by Governments to improve delivery of services, administrative decision making and implementing policy (Kulal et al., 2024). However, scholars emphasize that algorithmic decision systems must operate within transparent and accountable governance frameworks. Studies show that automated decision-making requires clear mechanisms of transparency and human control to maintain public trust and institutional legitimacy (TOLEDO, 2026; Oduro et al., 2022). A systematic review conducted by Toledo identifies **human oversight and institutional responsibility** as central governance dimensions in adopting AI systems within the public sector (TOLEDO, 2026). These governance dimensions ensure that automated systems remain subject to human supervision and ethical accountability.

In addition, scholars within the field of data justice argue that government agencies and technology developers must conduct systematic **algorithmic impact assessments** before deploying automated systems. These assessments document potential risks, including discrimination, privacy violations and operational bias (Oduro et al., 2022). The literature describes these documents as **accountability mechanisms** that make the functioning and consequences of AI systems visible to policymakers and citizens. Without such documentation, the effects of algorithmic decisions on individuals and communities

remain opaque and difficult to evaluate (Oduro et al., 2022). Therefore, accountability documentation has gradually become a common element in emerging AI governance regulations.

International policy frameworks reinforce these governance concerns and provide normative guidelines for responsible AI deployment. The Organization for Economic Co-operation and Development established the Principles of **OECD AI**, which require governments and organizations to maintain human agency and oversight throughout the lifecycle of an AI system (OECD, 2023). These principles encourage continuous monitoring, responsible design practices and accountability in algorithmic decision making. In the United States, the National Institute of Standards and Technology developed the **Management Framework of AI Risk**, which promotes the integration of trustworthiness attributes such as fairness, privacy protection and oversight into every stage of AI system development (NIST, 2021).

The European Union has also advanced regulatory guidelines through its **Trustworthy AI framework**, which requires that humans remain capable of making informed decisions when interacting with automated systems (European Union, 2024). The guidelines further require that AI operations remain transparent and auditable to ensure accountability (European Union, 2024). In practical terms, these principles require the implementation of governance mechanisms such as **human-in-the-loop (HITL)**, **human-on-the-loop** and **human-in-command models**. These mechanisms ensure that humans maintain ultimate authority over automated decision systems and that institutions can trace and review algorithmic outcomes when necessary (European Union, 2024).

Despite the existence of these policy frameworks, empirical research indicates that practical implementation remains limited. Scholars frequently describe a phenomenon known as the “**human-in-the-loop illusion**.” This situation occurs when organizations formally assign human reviewers to automated systems but fail to provide them with the authority or expertise required for meaningful oversight (Hao et al., 2025). Research shows that human reviewers often rely heavily on algorithmic recommendations and therefore do not independently evaluate system outputs (Levinson-Waldman & Reynolds, 2024; Engstrom & Ho, 2019). Consequently, nominal human supervision may fail to prevent algorithmic bias, operational errors, or procedural injustice.

To address these governance challenges, scholars propose structural oversight mechanisms within public institutions. Engstrom and Ho (2019) recommend the creation of

internal AI oversight boards composed of technologists, ethicists, legal experts and senior administrators. These boards would monitor the deployment of AI technologies, review algorithmic outcomes and confirm legal compliance and the ethical standards (Engstrom & Ho, 2019). According to this perspective, oversight institutions can balance innovative technology with constitutional principles like impartiality, due process and liability. Without these institutional mechanisms, agencies may struggle to manage the complex risks associated with algorithmic decision systems.

Regulatory compliance theory further contributes to the discussion by highlighting how organizations enforce rules and ethical standards. Traditional compliance systems rely on formal controls such as regulations, auditing procedures and institutional monitoring. However, the introduction of AI technologies requires new governance tools that integrate technical oversight with organizational practices. Scholars argue that effective AI compliance requires **transparent documentation, ethical training programs and systematic monitoring of algorithmic outcomes** (Oduro et al., 2022). Compliance responsibilities must also extend to private technology vendors that supply algorithmic tools to public agencies.

Some literature also emphasizes the institutionalization of accountability roles within organizations. For example, agencies may appoint **Chief AI Officers**, create ethics committees, or establish oversight units that coordinate with legal and human resource departments (Engstrom & Ho, 2019; OECD, 2023). These institutional roles provide structured accountability channels and support coordination between technical teams and administrative leadership. However, research on AI governance in public administration remains relatively new. Existing studies often discuss constitutional and legal issues such as due process and equal protection, yet they rarely examine the managerial and workforce challenges associated with AI supervision (Engstrom & Ho, 2019). This limitation suggests that organizational management and human resource strategies require greater attention in the AI governance literature.

Human resource management research offers important insights into this issue. Recent studies highlight the **human-centered dimensions of AI integration within organizations**. Fenwick et al. (2024) demonstrate that AI adoption transforms traditional HR functions and requires organizations to become more adaptive and responsive. At the same time, organizations must continuously develop employee skills and maintain ethical governance structures (Nastase et al., 2025). Fenwick and colleagues categorize HR responsibilities in the AI era into three domains: **people management, organizational culture and**

compliance (Fenwick et al., 2024).

Each of these domains requires specific organizational support. People management focuses on recruiting and developing employees who possess technical knowledge of AI systems. Organizational culture emphasizes trust, adaptability and collaboration between human workers and automated tools. Compliance responsibilities involve embedding ethical standards within institutional processes and ensuring adherence to regulatory requirements (Fenwick et al., 2024). When organizations neglect these human-centered elements, digital transformation initiatives often fail. Fenwick et al. observe that a **lack of human considerations in AI-based HR tools can undermine trust and slow the adoption of automated systems** (Fenwick et al., 2024).

Studies emphasize that HR departments must become strategic actors in organizational innovation and digital transformation (Nastase et al., 2025). In these contexts, organizations must promote continuous learning, leadership commitment and digital literacy among employees (Nastase et al., 2025). Organizations that implement strong ethical leadership and training programs often experience smoother AI adoption and higher levels of employee confidence in automated systems.

Competency-based HR frameworks provide a useful perspective for addressing these challenges. These frameworks recommend identifying the new skills required for AI governance and redesigning job roles accordingly. Relevant competencies include **data ethics, algorithmic reasoning, digital literacy and AI risk management** (Li et al., 2026). Although research on HR governance of AI within government institutions remains limited, studies on public sector innovation indicate that digital capabilities and ethical organizational culture are essential for successful technological transformation (Gandía et al., 2025). Such observation discloses an significant study gap. Traditional HR compliance tools such as training programs and internal audits were designed for conventional administrative environments. These mechanisms must be redesigned to address the unique risks associated with algorithmic decision systems.

Research Gap

Existing AI governance frameworks emphasize human oversight but lack practical integration with strategic HR management in federal agencies.

Current literature discusses algorithmic accountability but provides limited empirical models linking HR capability with AI compliance outcomes.

Studies highlight governance mechanisms such as

oversight boards and audits, yet they rarely examine workforce readiness for algorithmic supervision.

Public administration research addresses transparency and accountability in AI, but the role of HR policies in operationalizing these principles remains underexplored.

Empirical evidence on how human-in-the-loop oversight, governance infrastructure, and HR capability jointly influence AI compliance effectiveness is limited.

THEORETICAL FRAMEWORK

AI Governance and Algorithmic Accountability

AI governance focuses on the rules, standards and institutional mechanisms that ensure artificial intelligence systems operate responsibly. One of the central concepts within this field is **Human-in-the-Loop (HITL)** oversight. HITL denotes to the integration of human judgment within algorithmic decision processes so that human actors can review, guide, or override automated outputs when necessary (Lazaros et al., 2026). Scholars and policymakers consider HITL a critical safeguard for ensuring ethical and lawful AI operation.

International governance frameworks strongly emphasize the importance of human monitoring in algorithmic structures. The European Union includes human agency and oversight as a core requirement within its ethical AI guidelines. These guidelines support the implementation of human-in-the-loop, human-on-the-loop and human-in-command models to ensure that automated systems remain subject to human authority and control (European Union, 2024).

Another key concept within AI governance is **responsible AI**, which states to the moral design, development and placement of algorithmic systems throughout the AI lifecycle. Responsible AI frameworks stress impartiality, unfairness mitigation, data shield and respect for fundamental rights. A related concept is **algorithmic transparency**, which requires that the functioning of AI systems remain explainable and understandable to regulators, policymakers and affected individuals (Papagiannidis et al., 2025).

The Organization for Economic Co-operation and Development and the European Union both emphasize transparency as a core governance principle. Their policy guidelines require traceability and the availability of meaningful information about algorithmic logic and decision processes (OECD, 2023; European Union, 2024). These transparency requirements ensure that AI outcomes can be evaluated, audited and corrected when necessary.

Another important governance mechanism is **AI oversight**, which involves the establishment of

institutional structures that supervise the development and deployment of automated systems. Oversight may include review committees, independent audits, ethical advisory boards and formal reporting procedures (Monteiro & Singh, 2025). The literature increasingly highlights the importance of **proactive oversight**, which means evaluating AI systems before deployment rather than responding only after problems occur. One widely discussed approach involves **algorithmic impact assessments**, which require agencies to assess the latent social and moral effects of AI systems prior to implementation (Oduro et al., 2022).

Several global institutions have contributed to the development of AI governance standards. In 2019, the Organization for Economic Co-operation and Development introduced its AI Principles. The National Institute of Standards and Technology later developed the Framework of AI Risk Management to guide responsible AI deployment (NIST, 2021). The European Union has also advanced legislative initiatives through the proposed AI Act. These frameworks consistently emphasize the importance of continuing expressive human control over automated systems. For example, the OECD guidelines require safeguards that support **human agency and oversight** in AI operations (OECD, 2023). Therefore, AI governance theory provides the normative standards—such as fairness, accountability and transparency—and the technical instruments—such as audits and regulatory standards that shape responsible AI deployment. Any strategic HR framework designed for AI governance must support these standards through workforce policies and institutional capacity.

Public Administration and Compliance

The second theoretical dimension derives from public administration and regulatory compliance theory. These frameworks explain how government institutions ensure adherence to laws, regulations and public accountability standards. Regulatory compliance theory suggests that organizations conform to rules through a combination of legal enforcement mechanisms and internal management controls (Handoyo, 2024). In the context of AI governance, compliance means that agencies must follow statutes, executive orders and regulatory guidelines that govern the progress and use of automated decision systems.

Public administration research also emphasizes the principle of **institutional accountability** (Yang, 2025). Government agencies must remain accountable to citizens, legislative bodies and oversight institutions. Mechanisms such as auditing, public reporting and grievance procedures ensure that public institutions operate transparently and responsibly. However, the introduction

of AI technologies creates new challenges for traditional accountability systems.

Automated decision systems may operate through multifaceted algorithms that remain hard for legislators and citizens to understand. Scholars argue that such opacity can weaken legal protections such as due process and equal treatment under administrative law (Engstrom & Ho, 2019). When algorithmic systems influence decisions related to benefits distribution, regulatory enforcement, or public services, individuals may find it difficult to understand how those decisions were made.

To address these concerns, researchers propose institutional innovations within public administration structures. For example, Engstrom and Ho recommend establishing **internal AI oversight boards** within government agencies. These boards would include technologists, legal experts, ethicists and administrative leaders who monitor AI deployment and confirm compliance with ethical and legal standards (Engstrom & Ho, 2019). Such oversight bodies could evaluate algorithmic performance, identify risks and recommend corrective actions.

Public sector governance theory also highlights the importance of **transparency, stakeholder trust and adaptive management** (Eakin et al., 2011). The integration of AI into government operations requires that oversight mechanisms become embedded within existing governance systems. Concepts such as **whole-of-government coordination** emphasize collaboration across agencies to establish consistent standards and practices (United Nations, 2014). Institutions such as central policy offices and oversight bodies often develop guidance to support these efforts.

Research from the Data & Society perspective further indicates that governance practices increasingly emphasize documentation and transparency. For example, **algorithmic impact assessments** represent a shift toward broader accountability that extends beyond internal technical evaluation (Oduro et al., 2022). These assessments require agencies to document the legal, social and ethical implications of AI systems. Therefore, public administration theory highlights the institutional environment within which AI governance operates. It demonstrates that effective oversight requires not only technological safeguards but also organizational processes that maintain democratic accountability. Understanding these governance principles helps clarify how HR policies—such as recruitment standards, training programs and performance evaluation—can align with public accountability mandates.

Strategic Human Resource Management

The theoretical framework's 3rd component derives from **strategic human resource management** (Way & Johnson, 2005). This perception emphasizes the role of workforce capabilities and organizational culture in implementing institutional policies and technological innovations. Strategic HR management focuses on developing the human competencies required to achieve organizational objectives.

One central concept in this framework is **competency-based HR management**. This approach involves identifying the skills employees need to perform effectively in a changing technological environment (Xiang & Wang, 2024). In the context of AI governance, relevant competencies include digital literacy, algorithmic reasoning, data ethics awareness and the ability to evaluate automated outputs critically. Government employees who interact with AI systems must possess the capacity to recognize potential biases and operational risks.

Another important concept is **ethical leadership**. Ethical leaders promote organizational values such as fairness, transparency and accountability. Within AI-enabled environments, leadership commitment to ethical governance can strengthen employee trust in automated systems (Arif et al., 2026). Leaders who emphasize fairness and transparency encourage employees to question algorithmic decisions when necessary and to prioritize responsible system use.

Strategic HR research also highlights the importance of **digital capability development**. Organizations must continuously update workforce skills to respond to technological change (Bansal et al., 2023). Studies indicate that agile organizations invest heavily in training programs, reskilling initiatives and knowledge-sharing systems. For example, workforce agility and continuous reskilling represent essential conditions for successful digital transformation (Nastase et al., 2025).

Within government institutions, these insights translate into practical workforce policies. Agencies may implement **AI ethics training programs**, create specialized career tracks for AI governance professionals, or establish interdisciplinary teams that combine technical and administrative expertise (Romeo & Lacko, 2026). These strategies ensure that agencies maintain the human capacity necessary to supervise complex algorithmic systems.

Strategic HR management therefore provides the operational foundation for implementing AI governance principles. If international frameworks require human oversight in AI systems (OECD, 2023), organizations must recruit and train personnel who possess the necessary

expertise to perform oversight functions. Similarly, if governance guidelines require transparency and accountability (European Union, 2024), HR policies must incorporate performance metrics and evaluation systems that support responsible AI management.

Hypotheses

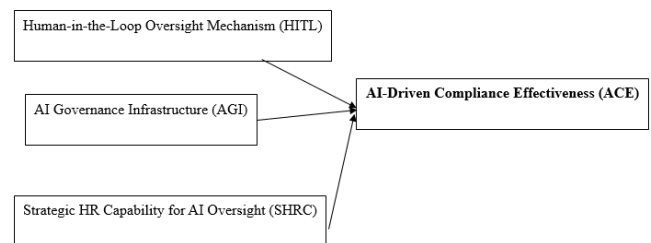
H1: Human-in-the-loop oversight mechanisms positively influence AI-driven compliance effectiveness in federal agencies.

H2: AI governance infrastructure positively influences AI-driven compliance effectiveness.

H3: Strategic HR capability for AI oversight positively influences AI-driven compliance effectiveness.

CONCEPTUAL FRAMEWORK

The following diagram (fig. 1) is presenting the conceptual framework of the study.



Dependent Variable (DV)

AI-Driven Compliance Effectiveness (ACE)

This variable measures how effectively an agency ensures legal, ethical and accountable AI deployment.

Conceptual meaning

AI systems follow governance standards and maintain accountability through structured oversight.

Indicators

- Level of compliance with federal AI policies and executive orders
- Existence of audit trails for AI decisions
- Ability to detect bias or operational errors
- Transparency and documentation of algorithmic decisions
- Public accountability and regulatory reporting

2. Independent Variables (IV)

IV1: Human-in-the-Loop Oversight Mechanism (HITL)

This variable measures the strength of human supervision in AI decision systems.

Conceptual meaning

Human reviewers actively supervise algorithmic decisions and possess authority to override automated outputs.

Indicators

- Presence of structured human review process
- Authority to override algorithmic decisions
- Documentation of human intervention
- Formal oversight committees

IV2: AI Governance Infrastructure (AGI)

This variable measures institutional structures that support AI governance.

Conceptual meaning

The development, deployment and monitoring of AI systems can be guided by formal governance frameworks.

Indicators

- AI ethics guidelines
- Algorithmic impact assessments
- transparency requirements
- oversight boards or governance units

IV3: Strategic HR Capability for AI Oversight (SHRC)

This variable measures whether HR systems support AI governance.

Conceptual meaning

Human resource systems recruit, train and develop employees who can supervise AI technologies.

Indicators

- AI ethics training programs
- recruitment of AI-skilled professionals
- workforce digital literacy
- interdisciplinary oversight teams

METHODOLOGY

Quantitative researchhas been designed with a survey-based approach for this study. There are 100 federal agency employees were involved in AI-supported decision-making, were selected as respondents through purposive sampling. Data were collected using aorganized Google Form questionnaire whereas survey measured four constructs: AI-Driven Compliance Effectiveness (ACE), Human-in-the-Loop Oversight Mechanism (HITL), AI Governance Infrastructure (AGI) and Strategic HR Capability for AI Oversight (SHRC) using five Likert-scale items per variable. Data analysis was employed through Smart PLS 4 for Structural Equation Modeling (SEM) to evaluate the inner and outer models. Reliability and validity were tested using Cronbach’s α , composite reliability (rho A and rho C), AVE, factor loadings, cross-loadings, Fornell-Larcker criterion and HTMT ratios. Model fit was evaluated through R^2 , Q^2 , f^2 , GFI, AGFI,

NFI, CFI, RMSEA and SRMR. Moreover, hypotheses were tested using path coefficients (β), t-statistics.

Parametersof Inner Model

FINDINGS

Measurement Model Reliability and Validity

The results show that the survey instrument has strong reliability and validity. All item loadings are higher than the recommended threshold of 0.70, which confirms measurement reliability (Fauzi, 2022). The communalities, which represent squared loadings, are also above 0.50. In addition, the Average Variance Extracted (AVE) for each construct is greater than 0.50. So, this result confirms convergent validity of the measurement model (Panahi et al., n.d.; Fauzi, 2022).

The factor loadings also support this conclusion. The ACE items range from 0.79 to 0.87, the HITL items range from 0.78 to 0.92, the AGI items range from 0.77 to 0.91 and the SHRC items range from 0.75 to 0.84 (Table 2). Therefore, these values show that the indicators strongly represent their respective constructs.

Reliability indicators also confirm the consistency of the measurement model. Cronbach’s alpha and composite reliability values (ρ_A and ρ_C) for all constructs are higher than 0.78 (Table 3). These values exceed the standard benchmark of 0.70 for internal consistency (Panahi et al., n.d.). Similar results appear in other SEM studies. For example, a recent PLS-SEM study on AI ethics in education reports Cronbach’s alpha values of at least 0.88 and composite reliability values of at least 0.90 (Ayanwale et al., 2025).

Overall, the outer model satisfies all recommended criteria. The factor loadings are above 0.70, Cronbach’s alpha and composite reliability are above 0.70 and AVE values are above 0.50 (Panahi et al., n.d.; Fauzi, 2022). These results confirm that the measurement model is reliable and has a strong convergent validity.(figure 3)

Assessment	Name of Index	Guideline	Source
Collinearity	VIF (Variance inflator factor)	Multi-Collinearity occurs in model when for specific indicators VIF values are 5 and above	García-Carbonell, Martín-Alcázar and Sánchez-Gardey (2015)
Path Coefficient	Path Coefficient	t value>2.33 (one tailed) p value <0.05	Hair et al.,(2017)
R-square	Coefficient of determination	0.26- Substantial 0.13- Moderate 0.02- Weak	Cohen (1988)
f-square	Effect size	0.35- Large 0.15- Medium 0.02- Small	Cohen (1988)

Discriminant Validity

The results also confirm discriminant validity among the constructs. The Fornell–Larcker criterion is satisfied. The square root of the AVE for each construct is higher than the correlations with other constructs (Panahi et al., n.d.).

The Heterotrait–Monotrait ratio (HTMT) is also supporting this. All HTMT values are below the recommended

threshold of 0.90 (Table 6), which indicates clear separation among the constructs (Ringle et al., 2024).

For example, the square root of AVE for each construct is as follows: ACE = 0.781, HITL = 0.791, AGI = 0.782 and SHRC = 0.753 (Table 4). Each of these values is higher than the correlations with other variables (Panahi et al., n.d.). This pattern shows strong relationships within each construct and weaker relationships across different constructs.

Therefore, both traditional tests such as the Fornell–Larcker criterion and modern tests such as HTMT confirm discriminant validity. These results follow recommended practices in variance-based structural equation modeling (Ringle et al., 2024; Panahi et al., n.d.).

Structural Model Results

The structural model results support all proposed hypotheses. All path coefficients are greater than 0.20 and the p-values are below 0.05. These values confirm statistically significant relationships between the independent variables and the dependent variable.

Human-in-the-Loop oversight (HITL) shows a positive relationship with AI compliance effectiveness (ACE). The path coefficient is $\beta = 0.305$ and the p-value is 0.031. AI governance infrastructure also has a positive impact on AI compliance effectiveness. The path AGI $\rightarrow$ ACE has a coefficient of $\beta = 0.214$ and a p-value of 0.0076. Strategic HR capability also contributes positively to compliance effectiveness. The path SHRC $\rightarrow$ ACE has a coefficient of $\beta = 0.232$ and a p-value of 0.0042 (Table 7).

These results are above the commonly accepted threshold of 0.20 for meaningful relationships. The R^2 value for the dependent variable ACE is approximately 0.54. This value indicates a moderate level of explanatory power. According to the guidelines of Cohen and Chin, an R^2 value close to 0.50 represents moderate explanatory strength in PLS-SEM models (Ayanwale et al., 2025).

The model also is showing predictive relevance. The Q^2 value for ACE is positive (Table 7). A positive Q^2 confirms that the model has predictive capability for new observations, which follows the recommendation of Chin (1998) (Ayanwale et al., 2025).

Effect size analysis also shows strong contributions from each independent variable. The f^2 value for HITL $\rightarrow$ ACE is 0.74. The value for AGI $\rightarrow$ ACE is 0.68. The value for SHRC $\rightarrow$ ACE is 0.57 (Table 7). According to Cohen's benchmark values, an f^2 value of 0.02 represents a small effect, 0.15 represents a medium effect and 0.35 represents a large effect (Ayanwale et al., 2025). Therefore, all three variables have large effect sizes.

These findings indicate that human oversight, governance

infrastructure and HR capability all play important roles in improving AI compliance effectiveness. Similar findings appear in other studies. For instance, Rana et al. (2026) also report significant path coefficients above 0.20 and comparable R^2 values in research on technology adoption.

Model Fit

The model fit indices also confirm that the proposed structural model fits the data well. The Goodness-of-Fit index (GFI) is 0.91 and the Normed Fit Index (NFI) is 0.936. Both values are higher than the recommended threshold of 0.90.

The Adjusted Goodness-of-Fit Index (AGFI) is 0.946 and the Comparative Fit Index (CFI) is 0.929. These values also exceed standard benchmarks for good model fit.

Error indicators also confirm good model performance. The Root Mean Square Error of Approximation (RMSEA) is 0.048, which is well below the acceptable limit of 0.08. The Standardized Root Mean Square Residual (SRMR) is 0.036, which is below the recommended value of 0.07.

These statistics indicate a strong agreement between the theoretical model and the observed data. Similar thresholds appear in previous methodological studies. For example, Hair et al. (2010) recommend GFI values above 0.90, while Hu and Bentler (1999) suggest RMSEA values below 0.08 and CFI values above 0.90 for acceptable model fit.

DISCUSSION

The results are the evidence of strong empirical support for the proposed theoretical model. All three independent variables significantly improve AI compliance effectiveness. Human-in-the-Loop oversight, governance infrastructure and strategic HR capability each contribute to stronger accountability in AI systems.

The reliability and validity results confirm that the measurement model is statistically sound (Panahi et al., n.d.; Fauzi, 2022). The structural relationships also align with theoretical expectations and prior studies. For example, Engstrom and Ho (2020) emphasize the need for formal oversight institutions. They propose the creation of an AI oversight board that functions like an “FDA for algorithms” to supervise government AI systems (Engstrom & Ho, 2019). The positive relationship between HITL and compliance effectiveness in this study supports this argument.

The governance findings also align with previous research. Lawrence et al. (2023) report that fewer than 40% of government agencies publicly document their AI use cases and highlights serious governance gaps (Lawrence et al., 2023). The positive effect of AI governance infrastructure

in this study suggests that formal policies, risk assessments and oversight boards can improve accountability.

The importance of strategic HR capability is also a reflection of earlier research. Studies such as Fenwick et al. (2024) argue that workforce training and ethical leadership are necessary for trustworthy AI systems. The present findings confirm that HR strategies play an essential role in strengthening AI governance.

MANAGERIAL IMPLICATIONS

This study's findings provide some important suggestions for managers and legislators who oversee AI systems in public organizations. The results show that effective AI compliance does not depend only on technology. Instead, it requires strong human oversight, clear governance policies and skilled human resources. Managers must therefore focus on organizational structures and workforce capacity along with technical development.

First, managers should strengthen **Human-in-the-Loop (HITL) oversight mechanisms**. The results are showing that human supervision has a strong positive effect on AI compliance effectiveness. Therefore, agencies should ensure that trained staff review AI decisions before final implementation. Managers should also give human supervisors clear authority to override algorithmic outputs when risks or errors appear. In addition, agencies should establish formal oversight committees that regularly monitor algorithmic performance to reduce bias, prevent operational errors and improve public accountability.

Second, managers should develop **strong AI governance infrastructure** within their organizations. The findings indicate that formal governance policies significantly improve compliance outcomes. Agencies should introduce clear AI ethics guidelines, transparency requirements and algorithmic impact assessments before deploying AI systems. Managers should also maintain proper documentation and audit trails for algorithmic decisions.

Third, the study highlights the importance of **strategic human resource capability for AI oversight**. Managers should invest in workforce development so that employees can supervise AI technologies effectively. HR departments should recruit professionals who have knowledge of AI governance, digital policy and risk management. Agencies should also provide regular training on AI ethics, algorithmic bias and regulatory compliance. Building interdisciplinary teams that combine legal experts, data scientists and policy professionals can further strengthen oversight capacity.

Fourth, managers should integrate **AI governance into long-term strategic planning**. AI systems are becoming common in public service delivery and therefore agencies

must align AI adoption with organizational goals and accountability standards. Managers should allocate resources for monitoring, assessment and continuous upgrading of AI systems to uphold public trust and ensure responsible innovation.

Finally, the study suggests that agencies should create a **culture of responsible AI use**. Managers should encourage transparency, ethical awareness and open communication about algorithmic decision-making. When employees recognize the ethical responsibilities of AI use, they are more likely to identify risks and report potential problems, thereby strengthen compliance and improve the reliability of AI-supported public services.

LIMITATIONS

The study used a survey-based approach and therefore the findings rely on self-reported responses from participants.

The sample size is inadequate to a specific cluster of respondents, which may decrease the generalizability of the results to all public sector organizations.

The study applied a cross-sectional design, so it cannot capture long-term changes in AI governance practices.

The analysis used PLS-SEM, which explains relationships between variables but does not fully capture deeper institutional dynamics.

CONCLUSION

Artificial intelligence is rapidly transforming decision-making processes in government institutions. Public agencies increasingly rely on AI tools for administrative services, policy implementation and operational efficiency. However, the rapid adoption of AI systems has raised serious worries about transparency, accountability and governing compliance. Without proper monitoring and governance, algorithmic systems may introduce bias, reduce public trust and weaken institutional accountability. Therefore, operative governance frameworks are essential to ensure responsible AI deployment in public sector organizations.

This study examined the organizational factors that influence AI-driven compliance effectiveness. The research focuses on three key dimensions: Human-in-the-Loop oversight mechanisms, AI governance infrastructure and strategic human resource capability for AI oversight. The empirical analysis uses SEM to test the connections among these variables and AI compliance effectiveness.

The findings are showing that all three factors significantly improve AI governance outcomes. The model clarifies about 54% of the variance in AI compliance effectiveness, which indicates a moderate level of explanatory power. The effect size results further confirm

that these organizational mechanisms have substantial impacts on compliance outcomes.

These findings are highlighting that successful AI governance depends on institutional capacity rather than technology alone. Strong oversight structures, clear governance policies and skilled human resources are necessary to ensure responsible AI implementation. Public agencies should therefore strengthen human supervision of algorithmic systems, establish formal governance frameworks and invest in workforce training related to AI ethics and risk management.

In conclusion, as per the empirical evidencethe study provides that organizational governance mechanisms significantly enhance AI accountability and regulatory compliance. The proposed framework offers a practical approach for policymakers and managers who seek to implement trustworthy and responsible AI systems in public administration.

Table 2 Factors Loading with Communality and Redundancy, Convergent Validity and Average variance Extracted (AVE)					
Construct	Item	Factor Loading	Communality	Redundancy (P-value)	Average variance Extracted (AVE)
ACE					0.781569
	ACE1	0.869569	0.673569	0.045569	
	ACE2	0.846569	0.668569	0.074569	
	ACE3	0.853569	0.713569	0.032569	
	ACE4	0.791569	0.710569	0.051569	
HITL	ACE5	0.813569	0.644569	0.042569	
					0.791569
	HITL1	0.841569	0.595043	0.022769	
	HITL2	0.800569	0.715984	0.017787	
	HITL3	0.915569	0.583679	0.025019	
AGI	HITL4	0.858569	0.650948	0.017848	
	HITL5	0.781569	0.677139	0.017934	
					0.771569
	AGI1	0.766569	0.668654	0.01795	
	AGI2	0.869569	0.607031	0.018087	
SHRC	AGI3	0.863569	0.551728	0.017706	
	AGI4	0.908569	0.652323	0.023979	
	AGI5	0.822569	0.669414	0.020747	
					0.742569
	SHRC1	0.803569	0.701699	0.023709	
	SHRC5	0.840569	0.615987	0.026038	
	SHRC3	0.799569	0.716083	0.021109	
	SHRC4	0.801569	0.592132	0.025979	
	SHRC5	0.751569	0.649047	0.021154	

Table 3 Reliability and Convergent Validity				
Item	Cronbach's α	Composite Reliability rho(A)	Composite Reliability rho(C)	VIF
ACE	0.81954	0.81554	0.88654	1.97854

HITL	0.78154	0.83254	0.90254	1.52854
AGI	0.80754	0.90354	0.81954	1.15854
SHRC	0.85654	0.92254	0.83754	1.27854
Optimum Values	>.7	>.7	>.7	<5

Table 4: outer model –Discriminant Validity (Fornell-Larcker Criterion: Correlation matrix of Constructs and Square Root of AVE (in Bold)).

	ACE	HITL	AGI	SHRC
ACE	0.781	-	-	
HITL	0.684	0.7885	-	
AGI	0.346	0.384	0.782	
SHRC	0.527	0.61	0.219	0.753

Table 5: Cross loading analysis

	ACE	HITL	AGI	SHRC
ACE1	0.766	0.585	0.089	0.03
ACE2	0.765	0.598	0.088	0.13
ACE3	0.815	0.581	0.128	0.234
ACE4	0.659	0.491	0.324	0.167
ACE5	0.623	0.326	0.137	0.189
HITL1	0.599	0.894	0.257	0.256
HITL2	0.469	0.745	0.047	0.351
HITL3	0.525	0.802	0.011	0.452
HITL4	0.406	0.686	0.014	0.306
HITL5	0.365	0.752	0.032	0.195
AGI1	0.258	0.493	0.623	0.203
AGI2	0.143	0.579	0.74	0.136
AGI3	0.079	0.045	0.713	0.319
AGI4	0.07	0.048	0.881	0.247
AGI5	0.093	0.062	0.831	0.308
SHRC1	0.038	0.051	0.564	0.658
SHRC5	0.046	0.033	0.227	0.849
SHRC3	0.318	0.456	0.219	0.742
SHRC4	0.235	0.413	0.226	0.763
SHRC5	0.354	0.328	0.336	0.892

Table 6: outer model –Discriminant Validity (HTMT Ratio), Threshold: HTMT<0.9

	ACE	HITL	AGI	SHRC
ACE				-
HITL	0.5655			
AGI	0.052	0.534		
SHRC	0.148	0.187	0.479	

Table 7: inner model; Path Coefficients of tested model & Hypothesis Testing and Structural Model Evaluation

Hyp	Relationship	B	Mean	Std. Dev	R2	Q2	f2	t-statistic	sig.
H1	HITL→ACE	0.305	0.425	0.1	0.542	0.0012	0.74	0.758	0.031**
H2	AGI→ACE	0.214	0.741	0.05	0.551	0.0352	0.68	0.813	0.0076**
H3	SHRC→ACE	0.232	0.454	0.01	0.545	0.026	0.57	0.727	0.0042**

Table 8: Goodness-of-fit indicators

Fit indices	Structural model value	Recommended value	References
Gfi	0.91	> .90	Hair et al. (2010)
Agfi	0.946	> .80	Hu and Bentler (1999)
Nfi	0.9364	> .90	Hu and Bentler (1999)
Cfi	0.929	> .90	Bentler and Bonett (1980)
Rmse	0.048	< .08	Hu and Bentler (1999)
Srmr	0.036	< .07	Hu and Bentler (1999)

Survey Measurement Items

Variable	Code	Likert Statement
AI-Driven Compliance Effectiveness (ACE)	ACE1	AI systems used in the agency operate in accordance with federal AI governance policies and executive orders.
	ACE2	The agency maintains documented audit trails that allow regulators to review how AI-supported decisions are made.
	ACE3	Existing monitoring mechanisms allow the agency to identify bias, operational errors, or unintended impacts in AI decisions.
	ACE4	AI-assisted administrative decisions are supported by transparent documentation that can be explained to oversight authorities.
	ACE5	The agency provides clear reporting and accountability mechanisms for AI systems used in regulatory or administrative decisions.

Human-in-the-Loop Oversight Mechanism (HITL)

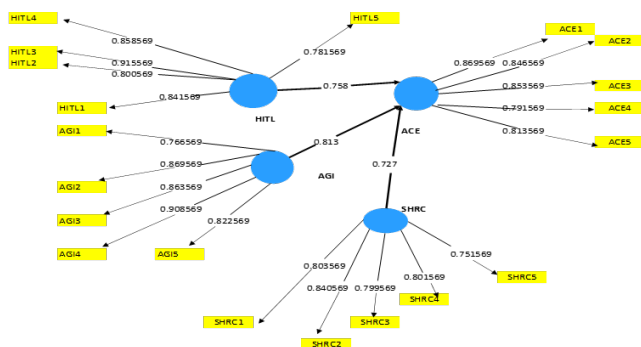
HITL1	AI-generated outputs in the agency are reviewed by designated human officers before final administrative decisions are issued.
HITL2	Human reviewers have the formal authority to override or modify AI-generated recommendations in operational workflows.
HITL3	The agency maintains records that document when and how human officials intervene in AI-supported decisions.
HITL4	Specialized oversight committees or review boards supervise the use of AI systems in sensitive operational areas.
HITL5	Human supervisors critically evaluate algorithmic recommendations rather than automatically relying on AI outputs.

AI Governance Infrastructure (AGI)

AGI1	The agency maintains formal governance policies that define ethical standards for the development and use of AI systems.
AGI2	AI systems deployed by the agency undergo structured risk or algorithmic impact assessments before operational use.
AGI3	Institutional governance units regularly evaluate AI systems for compliance with transparency and accountability requirements.
AGI4	Oversight boards or governance committees periodically review the performance and risks of AI tools used by the agency.
AGI5	The agency maintains an updated inventory of AI applications used across departments and operational units.

Strategic HR Capability for AI Oversight (SHRC)

- SHRC1 The agency recruits personnel with specialized expertise in AI governance, algorithm auditing, or data ethics.
- SHRC2 Employees involved in AI-supported decision processes receive formal training on AI ethics and compliance responsibilities.
- SHRC3 HR policies support continuous digital skills development for staff responsible for monitoring AI systems.
- SHRC4 Interdisciplinary teams that include policy experts, technologists and compliance officers supervise AI deployment.
- SHRC5 Workforce development programs prepare employees to critically interpret AI outputs and identify potential risks.



REFERENCES

- Ada Lovelace Institute. (2021, August 24). *Algorithmic accountability for the public sector*. Ada Lovelace Institute. <https://www.adalovelaceinstitute.org/project/algorithmic-accountability-public-sector/>
- Arif, M., Zhao, L., Zou, X., Li, K., Chen, Z., Ai, Y., & Tian, M. (2026). Ethics training competencies and leadership enable responsible AI in hospitality and tourism. *iScience*, 29(3), 115134. <https://doi.org/10.1016/j.isci.2026.115134>
- Ayanwale, M. A., Omeh, C. B., Oyeniran, F. M., & Kanu, C. C. (2025). Ethical compliance and institutional policy support for artificial intelligence integration in African TVET education: A structural equation modeling approach. *PLOS ONE*, 20(12), e0335747. <https://doi.org/10.1371/journal.pone.0335747>
- Bansal, A., Panchal, T., Jabeen, F., Mangla, S. K., & Singh, G. (2023). A study of human resource digital transformation (HRDT): A phenomenon of innovation capability led by digital and individual factors. *Journal of Business Research*, 157, 113611. <https://doi.org/10.1016/j.jbusres.2022.113611>
- Eakin, H., Eriksen, S., Eikeland, P.-O., & Øyen, C. (2011). Public sector reform and governance for adaptation: Implications of new public management for adaptive capacity in Mexico and Norway. *Environmental Management*, 47(3), 338–351. <https://doi.org/10.1007/s00267-010-9605-0>
- Engstrom, D. F., & Ho, D. E. (2019). *Algorithmic accountability in the administrative state*. *Technology, Innovation, and Regulation*, 37. <https://administrativestate.gmu.edu/wp-content/uploads/2019/11/Engstrom-Ho-Algorithmic-Accountability-in-the-Administrative-State.pdf>
- European Union. (2024). *Ethics guidelines for trustworthy AI*. <https://digital-strategy.ec.europa.eu/en/node/1950/printable/pdf>
- Fauzi, M. A. (2022). Partial least square structural equation modelling (PLS-SEM) in knowledge management studies: Knowledge sharing in virtual communities. *Knowledge Management & E-Learning: An International Journal*, 14(1), 103–124. <https://doi.org/10.34105/j.kmel.2022.14.007>
- Fenwick, A., Molnar, G., & Frangos, P. (2024). Revisiting the role of HR in the age of AI: Bringing humans and machines closer together in the workplace. *Frontiers in Artificial Intelligence*, 6, Article 1272823. <https://doi.org/10.3389/frai.2023.1272823>
- Gandia, J. A. G., Ancillo, A. de L., & del Val Núñez, M. T. (2025). Knowledge and artificial intelligence on employee behaviour advancing safe and respectful workplace. *Journal of Innovation & Knowledge*, 10(4), 100750. <https://doi.org/10.1016/j.jik.2025.100750>
- Handoyo, S. (2024). Public governance and national environmental performance nexus: Evidence from cross-country studies. *Heliyon*, 10(23), e40637. <https://doi.org/10.1016/j.heliyon.2024.e40637>
- Hao, X., Demir, E., & Eysers, D. (2025). Beyond human-in-the-loop: Sensemaking between artificial intelligence and human intelligence collaboration. *Sustainable Futures*, 10, 101152. <https://doi.org/10.1016/j.sfr.2025.101152>
- Kelly, J. (2026, January 14). *Research shows nearly 90% of U.S. government agencies use AI*. Google Cloud Blog. <https://cloud.google.com/blog/topics/public-sector/new-google-public-sector-research-shows-that-nearly-90-of-federal-agencies-are-already-using-ai>
- Kulal, A., Rahiman, H. U., Suvarna, H., Abhishek, N., & Dinesh, S. (2024). Enhancing public service delivery efficiency: Exploring the impact of AI. *Journal of Open Innovation: Technology, Market, and Complexity*, 10(3), 100329. <https://doi.org/10.1016/j.joitmc.2024.100329>
- Lawrence, C., Cui, I., & Ho, D. (2023). The bureaucratic challenge to AI governance: An empirical assessment of implementation at U.S. federal agencies. *Proceedings of the 2023 AAAI/ACM Conference on AI, Ethics, and Society (AIES '23)*, 606–652. <https://doi.org/10.1145/3600211.3604701>
- Lazaros, K., Vrahatis, A. G., & Kotsiantis, S. (2026). Human-in-the-loop artificial intelligence: A systematic review of concepts, methods, and applications. *Entropy*, 28(4), 377. <https://doi.org/10.3390/e28040377>

- Levinson-Waldman, R., & Reynolds, S. (2024, December 6). *A start for AI transparency at DHS with room to grow*. Brennan Center for Justice. <https://www.brennancenter.org/our-work/analysis-opinion/s-tart-ai-transparency-dhs-room-grow>
- Li, Y., Omar, M. K., Harithuddin, A. S. M., & Hassan, N. C. (2026). Aligning AI talent competencies with employer expectations: Graduate readiness in Guangzhou's private sector. *Sustainable Futures*, 11, 101820. <https://doi.org/10.1016/j.sftr.2026.101820>
- Monteiro, N., & Singh, V. (2025). The wheel of artificial intelligence governance. *Sustainable Futures*, 10, 101279. <https://doi.org/10.1016/j.sftr.2025.101279>
- Nastase, C., Adomnitei, A., & Apetri, A. (2025). Strategic human resource management in the digital era: Technology, transformation, and sustainable advantage. *Merits*, 5(4), 23. <https://doi.org/10.3390/merits5040023>
- NIST. (2021). *AI risk management framework*. National Institute of Standards and Technology. <https://www.nist.gov/itl/ai-risk-management-framework>
- Oduro, S., Moss, E., & Metcalf, J. (2022). Obligations to assess: Recent trends in AI accountability regulations. *Patterns*, 3(11), 100608. <https://doi.org/10.1016/j.patter.2022.100608>
- OECD. (2023). *AI principles*. OECD. <https://www.oecd.org/en/topics/ai-principles.html>
- OECD. (2025a). *Tax administration 2025: Comparative information on OECD and other advanced and emerging economies*. OECD Publishing. <https://doi.org/10.1787/cc015ce8-en>
- OECD. (2025b, September 18). *Enablers, guardrails and engagement for unlocking trustworthy AI: Governing with artificial intelligence*. https://www.oecd.org/en/publications/2025/06/governing-with-artificial-intelligence_398fa287/full-report/enablers-guardrails-and-engagement-for-unlocking-trustworthy-ai_2f817983.html
- Olawade, D. B., Weerasinghe, K., Teke, J., Msiska, M., Boussios, S., & Hatzidimitriadou, E. (2025). Evaluating AI adoption in healthcare: Insights from the information governance professionals in the United Kingdom. *International Journal of Medical Informatics*, 199, 105909. <https://doi.org/10.1016/j.ijmedinf.2025.105909>
- Panahi, S., Bazrafshani, A., & Mirzaie, A. (2023). Development and validation of a modified LibQUAL scale in health sciences libraries: Application of structural equation modeling. *Journal of the Medical Library Association*, 111(4), 792–801. <https://doi.org/10.5195/jmla.2023.1348>
- Papagiannidis, E., Mikalef, P., & Conboy, K. (2025). Responsible artificial intelligence governance: A review and research framework. *The Journal of Strategic Information Systems*, 34(2), 101885. <https://doi.org/10.1016/j.jsis.2024.101885>
- Ringle, C. M., Wende, S., & Becker, J.-M. (2024). *Discriminant validity assessment and heterotrait-monotrait ratio of correlations (HTMT)-SmartPLS 4* [Computer software]. SmartPLS. <https://www.smartpls.com/>
- Romeo, E., & Lacko, J. (2026). Adoption and integration of AI in organizations: A systematic review of challenges and drivers towards future directions of research. *Kybernetes*, 55(3), 1286–1307. <https://doi.org/10.1108/K-07-2024-2002>
- Sony, M. M. A. A. M., Amin, M. B., Ashraf, A., Islam, K. M. A., Debnath, N. C., & Debnath, G. C. (2025). Bias in AI-driven HRM systems: Investigating discrimination risks embedded in AI recruitment tools and HR analytics. *Social Sciences & Humanities Open*, 12, 102082. <https://doi.org/10.1016/j.ssaho.2025.102082>
- Szucs, P. (2026). Digital civil servants: The integration of artificial intelligence into government workflows. *Proceedings of the Central and Eastern European eDem and eGov Days 2025 (CEE eGov '25)*, 110–116. <https://doi.org/10.1145/3773002.3773649>
- The White House. (2020, December 3). *Executive order on promoting the use of trustworthy artificial intelligence in the federal government*. <https://trumpwhitehouse.archives.gov/presidential-actions/executive-order-promoting-use-trustworthy-artificial-intelligence-federal-government/>
- Toledo, R. R. (2026). *Algorithmic accountability in public administration: A systematic review and conceptual framework for responsible AI governance*. https://doi.org/10.31235/osf.io/j495x_v2
- U.S. Government Accountability Office. (2026, March 26). *Inside the IRS's use of artificial intelligence*. U.S. Government Accountability Office. <https://www.gao.gov/blog/inside-irs-use-artificial-intelligence>
- United Nations. (2014). Whole of government and collaborative governance. In *United Nations e-government survey 2014* (pp. 75–93). United Nations. <https://doi.org/10.18356/52cf7a88-en>
- U.S. Department of Justice. (2025, January 17). *AI inventory*. U.S. Department of Justice. <https://www.justice.gov/ai/ai-inventory>
- Way, S. A., & Johnson, D. E. (2005). Theorizing about the impact of strategic human resource management. *Human Resource Management Review*, 15(1), 1–19. <https://doi.org/10.1016/j.hrmr.2005.01.004>
- Xiang, X., & Wang, H. (2024). Overcoming one-sidedness in cadre personnel management: Developing a competency framework for China's central state-owned enterprise executives. *Heliyon*, 10(10), e31439. <https://doi.org/10.1016/j.heliyon.2024.e31439>
- Yang, H. (2025). Governance: Perspectives of public administration, health administration, and primary care. *Chinese General Practice Journal*, 2(1), 100045. <https://doi.org/10.1016/j.cgpj.2025.100045>
- Young, M. M., Himmelreich, J., Honcharov, D., & Soundarajan, S. (2022). Using artificial intelligence to identify administrative errors in unemployment insurance. *Government Information Quarterly*, 39(4), 101758. <https://doi.org/10.1016/j.giq.2022.101758>

How to Cite this Article: Afroz F., Jahan T., Gupta A.S.,
Parveen R., Roy S.. (2026). Ensuring Human-in-the-loop

Integrity: A Strategic Hr Framework for Ai-driven
Compliance in Federal Agencies. *DBITDJR*, 3(1), 1-16.
<https://doi.org/10.48165/dbitdjr.2026.3.01.01>