



UTILIZING ARTIFICIAL INTELLIGENCE AS A STRATEGIC RISK MANAGEMENT TOOL FOR PUBLIC SECTOR OPERATIONS AND AUDITING PROCESSES

KAMU SEKTÖRÜ FAALİYETLERİ VE DENETİM SÜREÇLERİ İÇİN STRATEJİK RISK YÖNETİM ARACI OLARAK YAPAY ZEKANIN KULLANILMASI

Oğuz Ümit TAMER¹

Bruce D. MCDONALD III²

Farouk HEMICI³

Georgia KONTOGEORGA⁴

ABSTRACT

Symbolizing a significant turning point in the historical landscape, AI is becoming an effective tool in today's public administration, not only for increasing capacity, quality, and speed in services, but also for strategic risk management. Regulators and algorithmic auditing play a central role in implementing fairness, transparency, and persistent controls against risks in AI systems. Discussing modern applications of AI, such as anomaly-based fraud detection, resource estimation, and continuous auditing, and their respective strengths and weaknesses, this study concludes that AI significantly enhances efficiency and oversight but also poses the risk of enshrining bias, opacity, and accountability gaps. By considering AI as a dual-use technology that demands a proactive paradigm of accountability, the study demonstrates that AI has the potential to help build more resilient and responsible public sector systems.

1- Principal Auditor, Turkish Court of Accounts, ORCID: 0000-0002-1848-9743.

2- Prof. Dr., Old Dominion University, USA, ORCID: 0000-0001-8963-8606.

3- Prof. Dr., University of Paris 1-Panthéon Sorbonne, ORCID: 0000-0002-2728-9372.

4- Dr., University of Paris 1-Panthéon Sorbonne, ORCID: 0000-0002-9830-324X.

Submitted/Gönderim: 28.10.2025, **Revised/Revizyon:** 05.12.2025, **Accepted/Kabul:** 05.12.2025

Corresponding/Sorumlu Yazar: Oğuz Ümit TAMER, oguz.tamer@sayistay.gov.tr

To Cite/Atıf: Tamer, O.Ü., McDonald, B.D., Hemici, F. and Kontogeorga, G. (2025). Utilizing Artificial Intelligence as a Strategic Risk Management Tool for Public Sector Operations and Auditing Processes. TCA Journal/Sayıştay Dergisi, 36(139), 659-680. DOI: <https://doi.org/10.52836/sayistay.1819606>

ÖZ

Tarihsel akışın önemli bir dönüm noktasını simgeleyen YZ, günümüz kamu yönetiminde sadece hizmetlerde bir kapasite, kalite ve hız artışının değil, aynı zamanda stratejik bir risk yönetiminin de etkili bir aracı haline gelmektedir. Düzenleyiciler ve algoritmik denetim, YZ sistemlerindeki risklere karşı adalet, şeffaflık ve kalıcı kontrollerin uygulanmasında merkezi bir rol oynamaktadır. Anomali tabanlı dolandırıcılık tespiti, kaynak tahmini ve sürekli denetim gibi yapay zekânın modern uygulamalarını ve bunların güçlü ve zayıf yönlerini tartışan bu çalışma, YZ'nin verimliliği ve denetimi büyük ölçüde artırdığı, ancak aynı zamanda önyargıyı yerleştirme, opaklık ve hesap verebilirlik boşlukları yaratma riskini de beraberinde getirdiği sonucuna varmaktadır. Yapay zekâyı proaktif bir hesap verebilirlik paradigması gerektiren çift kullanımlı bir teknoloji olarak ele alan çalışma, YZ'nin daha dayanıklı ve sorumlu kamu sektörü sistemleri oluşturmaya yardımcı olma potansiyeline sahip olduğunu göstermektedir.

Keywords: Artificial Intelligence, Public Sector, Risk Management, Auditing, Algorithmic Bias.

Anahtar Kelimeler: Yapay Zeka, Kamu Sektörü, Risk Yönetimi, Denetim, Algoritmik Önyargı.

INTRODUCTION

The extraordinary rapid development of the AI technology is making radical changes inevitable in the delivery of public services and the structure of public administration. The role of AI, which enables the intelligent interaction of information technologies with large data sets in strengthening decision-making and problem-solving capacity, is increasing daily and is driving revolutionary transformations in the functioning of states and societies. AI, which opens new horizons through the value provided by big data analytics and other new technological tools, has great potential to accelerate innovation and increase efficiency and effectiveness. These developments not only strengthen corporate governance in the public sector and support sustainability but also bolster strategic risk management and risk-based audit capabilities.

While AI reshapes democratic legitimacy mechanisms, it necessitates a redesign of the distribution and functioning of authority. This raises concerns that the adoption of AI in public services introduces systemic risks, specifically the potential for amplifying bias, increasing opacity, and weakening the implementation of democratic principles such as justice and equality. However, a more participatory, interactive, transparent, accountable, and ethically strengthened practice, far from weakening democratic governance, could make it much more effective and inclusive.

Algorithms developed in the auditing field through AI could pave the way for much more effective auditing by enabling continuous and timely analysis and instant reporting covering all areas. However, AI systems also pose significant risks, where inadequate oversight can conceal problems with algorithmic systems, potentially leading to the enshrinement of bias or legitimizing poorly designed systems. This raises concerns about the future of auditing, as institutions that fail to transform their assurance processes will inevitably provide false assurances about systems and statements, potentially obscuring accountability gaps (Goodman and Trehu, 2023: 302).

1. AI IN PUBLIC MANAGEMENT AND AUDITING

Artificial intelligence (AI) is increasingly portrayed as a transformative force in public-sector administration. Governments worldwide are piloting machine learning, natural language processing, and predictive analytics to enhance efficiency, reduce costs, and improve decision quality in areas such as welfare administration, budgeting, tax compliance, and fraud detection (Brookings Institution, 2021; Lee et al., 2024; Wirtz et al., 2019; Chung, 2025). Audit and oversight bodies are also experimenting with AI to strengthen assurance procedures and enable continuous monitoring that exceeds the limits of manual approaches (Issa, Sun, & Vasarhelyi, 2016; Kokina & Davenport, 2017; Yener et al., 2025). These trends suggest that AI is becoming not only a service-delivery tool but also a potential instrument of strategic risk management in government.

Yet the adoption of AI in sensitive public-sector contexts introduces significant risks. Algorithmic systems can amplify bias, obscure decision processes, and weaken public trust when deployed without adequate safeguards (Raji et al., 2020; Bandy, 2021). Recent controversies surrounding welfare fraud detection systems in Europe illustrate how poorly governed models can generate social harm (Wired, 2024). These cases underscore the need to align AI implementation with long-standing norms of public-sector governance such as fairness, proportionality, and due process (European Commission, 2024; OECD, 2019).

In response, regulators and standard-setting bodies have begun developing structured approaches to AI risk. The U.S. National Institute of Standards and Technology's AI Risk Management Framework emphasizes governance, measurement, and monitoring (NIST, 2023), while the European Union's Artificial Intelligence Act introduces a risk-based approach to high-impact public-sector applications (European Commission, 2024). Complementary guidance from organizations such as the OECD and G7 establishes principles intended to support the responsible adoption of AI in the public-sector AI (OECD, 2019; G7, 2023).

In practice, it is known that the use of AI in the audit sector is still in its infancy, and that both private sector audit institutions and supreme audit institutions (SAIs) in the public sector have conducted extensive research in this area. SAIs, driven by the need to use their resources effectively and efficiently, prefer to collaborate extensively in this area (Yener et al., 2025). These efforts are often conducted in collaboration with other SAIs because multiple perspectives are crucial for fully understanding the challenges of the task at hand. According to a survey conducted by the European Court of Auditors, SAIs will inevitably encounter some challenges in the medium term because of the use of AI systems by control institutions such as insufficient technical skills; compliance with legal, ethical, data protection and contractual obligations; the multidisciplinary complexity of the subject; the lack of business analysis skills in the institution; budgetary constraints etc. (European Court of Auditors, 2024). These findings underline the importance and logical basis of collaboration.

It is widely argued in the literature that AI can strengthen auditing processes in terms of efficiency, speed, risk prediction, and extensive data processing capacity. However, it can also increase vulnerabilities such as bias, lack of transparency, concentration of authority, and the erosion of procedural safeguards. Despite all its risks, algorithmic auditing, an inevitable necessity, is crucial to design it not merely as a means of capacity building or technological modernization, but as a key restructuring tool that will strengthen democratic governance in the public sector.

2. AI AND PROACTIVE RISK MANAGEMENT FRAMEWORK APPLICATIONS

AI systems are revolutionizing the way the public administration approaches risk, from a forensic and reactive mindset to a predictive and proactive one (Wirtz et al., 2019). This is true for both intrinsic administrative activities and for ancillary audit activities, in a context in which AI is employed not only for its effectiveness but also as a strategic means of anticipating financial, operating, and compliance risks.

2.1. AI in Core Public Sector Operations - Risk Mitigation

AI's strategic value in administrative and regulating authorities lies in filtering vast and intricate data streams to identify and categorize anomalies on a big scale. This function serves two valuable areas:

Traditional methods of detecting financial irregularities, such as welfare fraud and tax evasion, rely on existing business procedures and sporadic sampling, making them prone to significant false negatives and costly retrospective probes. By comparison, machine learning (ML) algorithms—notably supervised learning approaches—are trained on historic information to extract complex, non-evident patterns indicative of criminal activity (Brookings Institution, 2021). For example, predictive analytics may analyze tax returns or public benefit forms in real-time so agencies focus personnel on high-risk cases that would otherwise be auto-processed. This approach reshapes the role from a gatekeeping function to a sophisticated filtering of risk, thus minimizing leakage of public funds and optimizing compliance efforts (Wirtz et al., 2019).

In addition, AI tools make it possible, by the examination of information and data, to prevent misinformation and cyberattacks and to achieve greater transparency and competition, for example, in public procurement processes (Genaro-Moya et al., 2025). Increasing the traceability of processes through AI tools makes significant contributions to strengthening transparency and accountability (Yavuz, 2025), and ensuring transparency and accountability in public expenditures reduces the risk of corruption (Chen and Ganapati, 2023; Castro and Lopes, 2023; Köse and Polat, 2022).

Compared to traditional audit mechanisms, innovative approaches emerging with digitalization have the potential to make the anti-corruption process more effective and strengthen accountability by preventing asymmetric information. In this context, SAls must adapt their skills and techniques to be able to audit AI algorithms in their audits. Examples of good practices in this area are on the rise. Recent examples comprise; the use of ML/AI to detect non-existent schools claiming scholarship benefit (SAI India) or applying data mining and graph database to identify collusion issues in government procurement (SAI China) (SAI20, 2023).

However, risks of false positives and biased training data and consequent potential reputational and legal risks must not be overlooked. One such incident is the UK Post Office Horizon scandal, and while it predates modern generative AI, the underlying problem is the same:

Between 1999 and 2015, the Post Office prosecuted over 900 subpostmasters for theft and false accounting based entirely on financial discrepancies reported by the Horizon IT system (Post Office Horizon IT Inquiry, 2024). Auditors and investigators systematically accepted the system's output as absolute truth, ignoring evidence of software defects such as the "Dalmellington bug," which caused screen freezes that generated duplicate withdrawal entries when users repeatedly pressed keys (Bates v Post Office Ltd, 2019). Critically, the audit function failed to exercise professional skepticism, treating the "black box" data as a primary source of truth while disregarding corroborating human testimony or physical cash reconciliations. The High Court judgment noted that the system was not "robust" and that Fujitsu engineers retained undisclosed remote access to branch accounts, allowing them to alter transaction data without a visible audit trail—a direct violation of data integrity principles (Bates v Post Office Ltd, 2019). This reliance on flawed digital evidence without explainability or independent verification illustrates the "presumption of dependability" fallacy, where auditors default to trusting computer-generated records over human operators (Christie, 2020).

Aside from fraud prevention, AI also operates as a strategic early-warning system by increasing organizational foresight and operational resilience. Forecasting techniques, typically conducted using time-series analysis or natural language processing (NLP), assist governments in forecasting demand and making preparations to mitigate potential risks before they materialize. For

instance, ML systems are capable of predicting peaks in demand for a given social service or healthcare capacity (e.g., seasonal outbreaks of the influenza outbreaks or emergencies), and agencies are able to preposition staff and supplies (Wirtz, Weyerer, & Geyer, 2019).

An example of ML models being successfully deployed as dynamic predictive tools to manage capacity planning in high-pressure healthcare environments involves a pilot in England, UK, where the UK Health Security Agency developed ML pipelines to forecast short-term peak demand for health services during the 2022–23 winter season (Morbey et al., 2023). The models were trained on real-time syndromic surveillance data, such as telehealth cough calls and emergency department (ED) attendances for bronchiolitis in children, to forecast the timing and intensity of peaks for Respiratory Syncytial Virus (RSV) and Invasive Group A Streptococcal (iGAS) infections up to 28 days in advance. Crucially, the short-term forecasts demonstrated an ability to adapt and accurately predict a new, higher peak when the season became atypical due to an unexpected surge in iGAS cases, showcasing the superior adaptability of ML techniques over traditional time-series models in volatile scenarios. This proactive forecasting ability supports public health practitioners in optimizing critical resource allocation.

This application positions AI as a strategic operational risk management tool by increasing efficiency and reducing the possibility of service failure or crisis (Di Vaio et al., 2022). However, reliance on predictive models and therefore vulnerability to data quality issues or “black swan” events must also be kept in mind.⁵

2.2. AI in Public Sector Auditing Processes - Assurance and Oversight

The application of AI to public sector assurance processes greatly enhances the governance function through enhancements to both the breadth and the quality of assurance practices (Kokina & Davenport, 2017).

Public audits historically involved sampling on an infrequent basis. AI enables the Continuous Auditing (CA) model, where supervisory functions constantly examine the entirety of transactions and controls virtually in real-time (Issa, Sun, & Vasarhelyi, 2016; Coşkun and Bozkuş Kahyaoğlu, 2023). This

5- A black swan event is defined as an unpredictable, high-impact outlier that lies outside the domain of regular expectation and, despite its radical rarity, is rationalised and made to seem predictable only in retrospect.

is achieved largely by Robotic Process Automation (RPA) for the most mundane of data gathering and through a suite of ML algorithms for processing full data populations in a matter of instants for anomalies. This reduces drastically the built-in detection risk of conventional sampling and allows auditors to focus their expertise and judgment on those most troublesome or uncertain results, thus elevating the overall rigor of public accountability (Tiron-Tudor and Deliu, 2022).

AI techniques—particularly deep learning and sophisticated clustering algorithms—assist in recognizing subtle, multi-faceted anomalies that are likely to slip past human auditors or simple rule-based systems. Processing unstructured information (e.g., contracts and minutes of meetings) quickly using Natural Language Processing and matching it to structured financial data, AI enables a richer understanding of the activities of an agency. This enables not only recognizing possible fraudulent activities but also realizing important details about weaknesses in controls, thus boosting the capacity of the auditor to consider intrinsic risk and to create improved assurance plans (Kokina & Davenport, 2017). In the long run, the integration of AI tools is supposed to assist, not replace, human judgment, leading to a perceptible improvement in the quality of audits and rising confidence levels among stakeholders. Automation, analytics, and AI should be regarded as enablers akin to traditional computing technologies. These technologies are not intended to replace the human auditor; instead, they are expected to transform audit processes and the auditor's role (Köse and Polat, 2021; SAI20, 2023).

Nevertheless, as the development and implementation of AI systems by public administrations grow exponentially due to their potential to improve public services and reduce costs, new challenges and risks have also emerged. AI offers unprecedented opportunities to analyze vast quantities of data with speed and accuracy, but it also introduces new complexities in terms of governance, bias, and ethical use (Genaro-Moya et al., 2025).

3. RISKS OF AI, GAPS IN ACCOUNTABILITY, AND THE NEED FOR ALGORITHMIC AUDITING

The competitive advantage afforded to AI in the field of risk management is tempered with important responsibility and moral issues, especially when systems of this sort operate within essential public situations. This section focuses on the contradictions of AI by clarifying the ultimate causes of algorithmic risk and documenting the evolving governance structures necessary to bridge abstract responsibility and practical application.

3.1. Algorithmic Opacity, Bias, and Public Systemic Harm

AI governance's root problem is the intrinsic ambiguity and hidden potential for discrimination within sophisticated machine learning systems.

AI models, especially those employing deep learning techniques in domains such as classification and predictive scoring, function as opaque entities, rendering the underlying decision-making processes unintelligible to human operators and auditors (Bandy, 2021). Furthermore, much of the processing, storage, and use of information is performed by algorithms and in a non-transparent way, within a "black box" of virtually inscrutable processing, the content of which is unknown even to its programmers (Criado, 2021).

When AI systems are developed using historical data that reflects pre-existing systemic societal disparities (for instance, in areas like policing, lending, or welfare provision), they risk becoming mechanisms for perpetuating and amplifying discrimination. A notable instance of this can be observed in the implementation of welfare fraud detection systems across various European nations. The algorithms utilized to detect "anomalous" living patterns or "foreign affiliation" within Denmark's welfare framework have been demonstrated to disproportionately target minority populations and economically disadvantaged individuals for invasive investigations, thereby fostering an environment characterized by extensive surveillance and directing attention towards those whom the system is intended to assist (Wired, 2024). This scenario redefines a tool that was designed for the mitigation of financial risk into one that generates systemic social risks.

In a world of AI, risks of operations move from human errors to algorithm failures and data quality shortcomings. When an algorithm provides a negative outcome, an affected person faces an unbearable onus of proof and is unable

to challenge a decision of whose logic is blurred by proprietary techniques or technological sophistication. This lacuna of responsibility is also compounded by a lack of systematic documentation of an AI model's development, testing, and proposed use, thus making retrospective inquiry difficult or impossible (Raji et al., 2020).

3.2. The Response of Governance

Institutionalizing AI risk management to mitigate such risks, governments and international institutions have attempted to institutionalize responsibility through establishment of rigorous governance structures demanding risk analysis, transparency, and human oversight.

The European Union's Artificial Intelligence Act (AI Act) is a prime example of regulation based on a risk-based assessment, categorizing applications of AI on the basis of their potential to cause harm. Applications used for the purpose of public access, such as systems for credit scoring, or those used for detecting fraud and crime (especially in financial services and law enforcement), are designated as 'High-Risk' and thus trigger a comprehensive set of mandatory obligations (European Commission, 2024). These compulsory measures include establishing robust Quality Management Systems (QMS), ensuring the creation of detailed logging and traceability capabilities for audit trails, mandating rigorous human oversight mechanisms, and certifying high standards of data integrity and accuracy (European Commission, 2024). Crucially, the AI Act requires 'High-Risk' systems to undergo a conformity assessment, typically requiring third-party checks by a Notified Body, which functions as a regulatory audit to verify compliance before the systems can be placed on the market. In a similar vein, the OECD Guidelines on Artificial Intelligence (2019) also categorically refer to AI systems' responsibility, fairness, and transparency, and thus advocate for the creation of express mechanisms of redress, ensuring individuals have recourse against adverse, AI-driven decisions.

In the USA, the NIST Artificial Intelligence Risk Management Framework (AI RMF 1.0) is a voluntary and non-regulatory framework designed for organizations to manage AI-related risks across the entire lifecycle (NIST, 2023). It assists organizations in conceptualizing and handling AI risk through four fundamental functions: Govern, Map, Measure, and Manage. The Govern

function establishes organizational priorities and risk tolerance; Map identifies specific risks within the AI context (e.g., sources of bias or false positives); Measure involves developing quantitative metrics and testing methodologies for trustworthiness attributes like reliability and fairness; and Manage requires taking action on the measured risks and communicating them to stakeholders (NIST, 2023). While voluntary, the AI RMF 1.0 is becoming a de facto standard for best practice, influencing contractual obligations and supply-chain risk management across the federal sector and private industry by providing the necessary policy infrastructure to transform aspirational values—such as those articulated by the OECD—into tangible operational procedures for responsible AI use.

3.3. Algorithmic Auditing

Algorithmic auditing functions as a concrete method to implement these governance frameworks and address the accountability deficiencies highlighted by AI systems. This process encompasses a comprehensive assurance methodology that assesses an AI system's adherence to predefined principles and regulations (Raji et al., 2020). Audits are of two types, the first being the deploying agency's internal ones and the second being conducted using independent oversight groups. Both of them include systematic thinking throughout the AI lifecycle:

- During the design & training phase, the main focus is data quality & bias, and the main concern is whether the data is fair and representative. Indeed, the development of algorithmic auditing tools can exhibit critical data quality and bias problems during the initial design phase, before any model training even occurs. A potent real-life illustration is the Dutch System Risk Indicator (SyRI), an algorithmic tool designed by the government to detect potential welfare and tax fraud. The system's design phase incorporated an inherent bias by intentionally selecting proxy variables—including ethnicity, low income, and residence in specific postal codes—that are known to correlate disproportionately with protected characteristics, effectively baking systematic discrimination into the risk scoring mechanism (LVV v The Netherlands, 2020). Simultaneously, severe data quality challenges arose from the aggregation of seventeen different government registers (e.g., tax, employment, housing, debt), which were not designed for integrated risk analysis. This

combining of siloed, disparate data sets compromised the integrity and context of the input, making the final risk score non-transparent and virtually impossible to audit or verify. The resulting algorithmic audit function, which flagged thousands of citizens for intensive manual investigation based on these biased inputs, was ultimately ruled by the Hague District Court to violate the right to privacy and constituted a disproportionate risk of discrimination under the European Convention on Human Rights.

- During development, the main audit focus is model documentation and accountability is established through mandating the use of tools like Model Cards to document intended use, ethical considerations and performance metrics (Mitchell et al., 2019).
- During deployment and monitoring, the main audit focus is systemic impact and accountability mechanism is probing the system for unintended negative consequences, especially discriminatory outcomes, through continuous monitoring (Bandy, 2021).

By institutionalizing algorithmic auditing, the public sector can ensure its powerful AI technologies remain aligned with the values of society and thus converts the proactive risk detection framework into a Proactive Accountability Paradigm. This creates an essential feedback mechanism for responsible AI application, whereby risks intuited through auditing are fed back to improve and refine governance structures in place.

4. MERGING RISK MANAGEMENT AND AUDITING THROUGHOUT THE AI LIFECYCLE

Successful public sector AI incorporation requires a shift from a reactive, incident-based framework to an overarching Proactive Accountability Paradigm, which integrates (in a systematic fashion) the use of risk management and algorithmic assurance across the full AI lifecycle. This ushering in enables a recurring feedback mechanism, ensuring that the instruments used to identify the risk are continuously under strong, “always-on” review.

4.1. The Comprehensive Governance Framework: From Inception to Decommissioning

The linchpin of the synthesis is to treat AI, rather than as a static IT acquisition, as a dynamic system controlled by a formal, multi-phased lifecycle, similar to constructs such as the NIST AI RMF. “Bake-in” responsibility, rather than “bolt-on” after-deployment, must occur “baked-in” during the first phase (Raji et al., 2020).

Prior to training or fielding, the public agency should be required to record the purpose and system limits, as well as the ethical risk landscape (European Commission, 2024). This necessitates the required use of artifacts like Model Cards⁶ (Mitchell et al., 2019) and Datasheets for Datasets⁷ (Gebru et al., 2021) to record the provenance, technical description, as well as the performance metrics, most notably disparate impact across subpopulations. Auditability should be a prime procurement factor, such that the AI is “designed to be auditable.”

As the system is implemented, it falls into a stage that consists of Continuous Auditing of AI (CAAI). As this approach shifts the assurance process from the annual snapshot review to a nearly real-time, automated verification process (Minkinen, 2022). Transaction data, as also the AI system itself, is then monitored continuously based on specified metrics relating to bias, fairness, and accuracy, thereby utilizing AI to provide assurance on the AI itself (Erasmus & Kahyaoglu, 2024).

4.2. Interconnectedness between Continuous Monitoring and Continuous Auditing

Integrated Model gives rise to a significant connection between two central technological competencies: Continuous Monitoring and Continuous Auditing.

6- Model Cards (Mitchell et al., 2019) are standardized short documents designed to accompany trained machine learning models, detailing their performance characteristics, intended use cases, and limitations. They are crucial for transparent model reporting, especially by providing disaggregated performance metrics across different demographic or phenotypic groups to surface potential bias.

7- Datasheets for Datasets (Gebru et al., 2021) is a structured documentation framework—analogue to datasheets for electronic components—those mandates recording a dataset’s creation, composition, collection process, recommended uses, and potential ethical or legal concerns. This protocol enhances data provenance and accountability before model training.

An AI system embedded in a public sector (e.g., identification of fraud) is always monitoring transactional data both for anomalies as well as non-compliance, so as to regulate operational risk. These systems, capable of analyzing massive, complex, and multi-source datasets, will create a new audit culture that will replace traditional audit practices, necessitating a revamped auditor profile accordingly. Audits will be conducted largely with technology support, but auditor judgment and skepticism will continue to be the foundation of successful audit outcomes (Köse, 2023).

The audit process uses an autonomous AI layer to continuously monitor the effectiveness and procedural soundness of the continuous monitoring process. In the event the fraud detection model suffers a sudden loss of accuracy, a shift in predictive outcomes, or a bias to unjustly single out a specific protected group—more traditionally called algorithmic drift—the Continuous Auditing process generates an alert for human review (Vasarhelyi et al., 2018).

This integration ensures that the Proactive Risk Management Paradigm (AI detecting public sector risk) is reliably checked by the Proactive Accountability Paradigm (AI detecting AI risk), creating a robust system of algorithmic checks and balances.

CONCLUSION: PROSPECTS FOR PUBLIC SECTOR ASSURANCE IN THE ERA OF AI

The application of AI within the public sector reveals a significant duality: it acts as both an exceptional catalyst for enhancing efficiency and facilitating proactive risk identification, while concurrently serving as a substantial magnifier of systemic risks associated with bias and accountability.

This study contends that effectively managing this duality necessitates the establishment of a Proactive Accountability Paradigm, which should be founded on the compulsory incorporation of algorithmic auditing and risk governance throughout the comprehensive lifecycle of AI systems.

Insights obtained based on this review that are important to policymakers as well as assurance professionals include:

First, the governance must come before the deployment. While regulations like the EU AI Act and frameworks like the NIST AI RMF provide necessary ground infrastructure, the future of public assurance goes beyond the act of auditing AI systems to the assessment of the framework and organizational processes that govern them, which requires explicit compliance with principles of fairness, transparency, and human oversight.

Secondly, assurance must be ongoing. Conventional, periodic auditing lacks the structural capacity to align with the dynamic and evolving characteristics of machine learning models. The principle of Continuous Auditing of AI (CAAI) ought to be established as the norm for high-risk applications in the public sector, employing automation to guarantee real-time monitoring and intervention.

Third, the working force will undergo augmentation, not full automation. While AI will be responsible for the high-volume, high-repetition tasks related to data analysis and anomaly identification, the role of the human auditor will shift to strategic judgments—to evaluate the ethical implications of algorithmic decisions, to devise fairness judgments, to interact with affected citizens, and to consult on the AI deployment implications of public value (Kokina & Davenport, 2017). We face, as the future challenge to the public sector leadership, the need to quickly build the necessary digital and ethical competencies among the current workforce to operate efficaciously in this augmented environment (Wirtz, Weyerer, & Geyer, 2019).

In the long run, the intention is to shift the public sector's interaction with AI from a serendipitous jump into technological dependency to a routine practice of technological stewardship. By embracing routine, lifecycle-based review, government agencies would be empowered to tap the vast potential of AI to advance the common good.

It is critical to solidify the conceptual foundation of the Proactive Accountability Paradigm by addressing key research gaps in measurement and assurance. While emerging policies stress AI trustworthiness, standardized, domain-agnostic methods are needed to quantify attributes like fairness and non-discrimination. Future work should define and validate quantitative fairness metrics that go beyond demographic parity to capture disparate impacts in complex public services (as seen in the Dutch SyRI welfare-

fraud case). In parallel, next-generation continuous auditing tools must be developed. These should integrate real-time explainability (XAI - explainable artificial intelligence) features to make opaque model decisions transparent and to detect algorithmic drift – the gradual degradation of a model's integrity, accuracy, or fairness over time. If necessary, advances in auditable XAI and verifiable integrity metrics fail to materialize, and continuous assurance risks becoming a mere technical formality rather than a substantive safeguard for public-sector AI.

Future policy should shift from purely regulatory prescriptions toward practical implementation and capacity-building. A priority is establishing accredited, interdisciplinary training programs to bridge competency gaps among auditors and administrators. This will require structured cooperation between SAs and academic institutions to develop curricula in digital ethics, computational methods, and public law, thereby enhancing the workforce's ability for effective strategic judgment. Further research should clarify legal and administrative remediation mechanisms for adverse AI-induced decisions and focus on methods that will reduce the excessive burden of proof on affected citizens. Ultimately, an agenda addressing these challenges will help ensure that the adoption of AI evolves into sustained technological stewardship, helping AI serve the public good with resilience and responsibility.

REFERENCES

- Bandy, J. (2021). Problematic machine behavior: A systematic literature review of algorithm audits. *Proceedings of the ACM on Human-Computer Interaction*, 5(CSCW1), pp. 1–34. doi:10.1145/3449148.
- Bates v Post Office Ltd (2019). Bates v Post Office Ltd (No 6: Horizon Issues) [2019] EWHC 3408 (QB) [online]. <https://www.judiciary.uk/judgments/bates-others-v-post-office-ltd-judgment-no-6-horizon-issues/> (Accessed: 4 December 2025).
- Brookings Institutions – West, D.M. (2021). Using artificial intelligence and machine learning to reduce government fraud. *Brookings*. <https://www.brookings.edu/articles/using-ai-and-machine-learning-to-reduce-government-fraud/> (Accessed: 8 September 2025).
- Castro, C. and Lopes, I. C. (2023). E-government as A Tool in Controlling Corruption. *International Journal of Public Administration*, 46(16), 1137-1150.
- Chen, C. and Ganapati, S. (2023). Do Transparency Mechanisms Reduce Government Corruption? A Meta-Analysis. *International Review of Administrative Sciences*, 89(1), 257-272.
- Christie, J. (2020). The Post Office Horizon IT scandal and the presumption of the dependability of computer evidence. *Digital Evidence and Electronic Signature Law Review*, 17, pp. 49–70.
- Chung, H., Kara, B., McShea, M.F., Pathak, R. and Williams, D. (2025). Using Large Language Models to Forecast Local Government Revenue. *Public Finance Journal*, 2(2), 85–98. <https://doi.org/10.59469/pfj.2025.46>.
- Criado, J.I. (2021). Inteligencia Artificial (y Administración Pública). *Economía. Revista en Cultura de la Legalidad*, 20, pp. 348–372. doi:10.20318/economia.2021.6097.
- Coşkun, E. and Bozkuş Kahyaoğlu, S. (2023). Digital Transformation and The Role of Digital Technologies in The Transformation of Audit: Opportunities and Threats. (Eds. Önder, M. and Köse, H.Ö.) *Kamu Yönetiminde Denetim: Temel Paradigmalar, Değişim ve Yeni Yönelişler*. Ankara: Sayıştay Yayınları.
- Di Vaio, A., Hassan, R. and Alavoine, C. (2022). Data intelligence and analytics: A bibliometric analysis of human–artificial intelligence in public sector decision-making effectiveness. *Technological Forecasting and Social Change*, 174. doi:10.1016/j.techfore.2021.121201.
- Erasmus, L.J. and Kahyaoğlu, S.B. (eds.) (2024). *Continuous Auditing with AI in the Public Sector*. 1st edn. Boca Raton, FL: CRC Press. doi:10.1201/9781003382706.
- European Commission (2024). Artificial Intelligence Act: Regulation (EU) 2024/1689. Brussels: European Commission. <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A32024R1689> (Accessed: 20 September 2025).

- European Court of Auditors (2024). Special report 08/2024: EU Artificial intelligence ambition: Stronger governance and increased, more focused investment essential going forward. Publications Office of the European Union. <http://www.eca.europa.eu/en/publications/SR-2024-08> (Accessed: 8 November 2025).
- G7 (2023). G7 Toolkit for Artificial Intelligence in the Public Sector. Paris: OECD Publishing. <https://doi.org/10.1787/421c1244-en> (Accessed: 23 September 2025).
- Genaro-Moya, D., López-Hernández, A.M. & Godz, M. (2025). Artificial Intelligence and Public Sector Auditing: Challenges and Opportunities for Supreme Audit Institutions. *World*, 6(2), pp. 1–19. doi:10.3390/world6020078.
- Gebru, T., Morgenstern, J., Vecchione, B., Vaughan, J.W., Wallach, H., Daumé III, H. and Crawford, K. (2021). Datasheets for datasets. *Communications of the ACM*, 64(12), pp. 86–92. doi:10.1145/3458723.
- Goodman, E. P. ve Trehu, J. (2023). Algorithmic Auditing: Chasing AI Accountability. *Santa Clara High Technology Law Journal*, 39(3), 289–337.
- Issa, H., Sun, T. and Vasarhelyi, M.A. (2016). Research ideas for artificial intelligence in auditing: The formalization of audit and workforce supplementation. *Journal of Emerging Technologies in Accounting*, 13(2), pp. 1–20. doi:10.2308/jeta-10511.
- Kokina, J. and Davenport, T.H. (2017). The emergence of artificial intelligence: How automation is changing auditing. *Journal of Emerging Technologies in Accounting*, 14(1), 115–122. doi:10.2308/jeta-51730.
- Köse, H. Ö. (2023). Kamu Yönetiminde Denetimin İşlevleri, Dinamikleri ve Geleceği. (Eds. Önder, M. and Köse, H.Ö.) *Kamu Yönetiminde Denetim: Temel Paradigmalar, Değişim ve Yeni Yönelişler*, pp. 37–67. Ankara: Sayıştay Yayınları.
- Köse, H. Ö. and Polat, N. (2021). Dijital Dönüşüm ve Denetimin Geleceğine Etkisi, *Sayıştay Dergisi*, 32(123), 9–41.
- Lee, M.E.M., Hayes, D. and Maher, C.S. (2024). AI as a Budgeting Tool: Panacea or Pandora's Box? *Public Finance Journal*, 1(1), 49–65. <https://doi.org/10.59469/pfj.2024.6>.
- LVV v The Netherlands (2020). Landelijke Vereniging van Vertrouwenspersonen (LVV) et al. v The Netherlands. *Rechtbank Den Haag*, ECLI:NL:RBDHA:2020:1878 (05.02.2020). <https://uitspraken.rechtspraak.nl/details?id=ECLI:NL:RBDHA:2020:1878> (Accessed: 04.12.2025).
- Minkkinen, M., Laine, J. and Mäntymäki, M. (2022). Continuous auditing of artificial intelligence: A conceptualization and assessment of tools and frameworks. *Digital Society*, 1, article 21. doi:10.1007/s44206-022-00022-2.
- Mitchell, M., Wu, S., Zaldivar, A., Barnes, P., Vasserman, L., Hutchinson, B., Spitzer, E., Raji, I.D. and Gebru, T. (2019). Model cards for model reporting. in *Proceedings of the Conference on Fairness, Accountability, and Transparency (FAT19)*. Atlanta, GA: ACM, pp. 220–229. doi:10.1145/3287560.3287596.

- Morbey, R.A., Elliot, A.J., Harcourt, S. and Smith, G.E. (2023). Using machine learning to predict peak healthcare demand: a time-series analysis of syndromic surveillance data. *PLOS ONE*, 18(10), e0291932. <https://journals.plos.org/plosone/article?id=10.1371/journal.pone.0291932> (Accessed: 4 December 2025).
- NIST (National Institute of Standards and Technology) (2023). *AI Risk Management Framework (AI RMF 1.0)*. Gaithersburg, MD: U.S. Department of Commerce. <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf> (Accessed: 28.09.2025).
- OECD (Organisation for Economic Co-operation and Development) (2019). *OECD Principles on Artificial Intelligence*. Paris: OECD Publishing. <https://oecd.ai/en/ai-principles> (Accessed: 28 September 2025).
- Post Office Horizon IT Inquiry (2024). *Post Office Horizon IT Inquiry: Publications* [online]. <https://www.postofficehorizoninquiry.org.uk/> (Accessed: 04.12.2025).
- Raji, I.D., Smart, A., White, R.N., Mitchell, M., Gebru, T., Hutchinson, B., Smith-Loud, J., Theron, D. and Barnes, P. (2020). Closing the AI accountability gap: Defining an end-to-end framework for internal algorithmic auditing. in *Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency (FAT20) **. Barcelona, Spain: ACM, pp. 33–44. doi:10.1145/3351095.3372873.
- SAI20 (2023). *Compendium on Responsible Artificial Intelligence*. Supreme Audit Institution of India (SAI India). <https://sai20.org/storage/app/uploads/public/679/c97/4b0/679c974b0954d084882913.pdf> (Accessed: 20.11.2025).
- Tiron-Tudor, A. and Deliu, D. (2022). Reflections on the human-algorithm complex duality: Perspectives in the auditing process. *Qualitative Research in Accounting & Management*, 19(3), pp. 255–285. doi:10.1108/QRAM-04-2021-0059.
- Vasarhelyi, M.A., Alles, M.G. and Kogan, A. (2018). Principles of analytic monitoring for continuous assurance. in Vasarhelyi, M.A. and Kogan, A. (eds.) *Continuous Auditing: Theory and Application*. Bingley: Emerald Publishing, pp. 191–217.
- Wirtz, B.W., Weyerer, J.C. and Geyer, C. (2019). Artificial intelligence and the public sector—Applications and challenges. *International Journal of Public Administration*, 42(7), pp. 596–615. doi:10.1080/01900692.2018.1498103.
- Wired (2024). *Denmark’s welfare algorithms under fire*. Wired. <https://www.wired.com/story/denmark-welfare-algorithms-under-fire/> (Accessed: 28.09.2025).
- Yavuz, E. (2025). Digitalization in the Fight Against Corruption: Feasibility of a Blockchain-Based System in Türkiye. *TCA Journal/Sayıştay Dergisi*, 36 (137), 255-287. <https://doi.org/10.52836/sayistay.1675048>
- Yener, M., Charoenpol, M., Suntharanurak, S. and Köse, H.Ö. (2025). Strategic Cooperation of Supreme Audit Institutions of Thailand and Türkiye for Digital Transformation and Innovation in Public Sector Auditing. *TCA Journal/Sayıştay Dergisi*, 36 (136), 9-34, doi: 10.52836/sayistay.1633666

KAMU SEKTÖRÜ FAALİYETLERİ VE DENETİM SÜREÇLERİ İÇİN STRATEJİK RİSK YÖNETİM ARACI OLARAK YAPAY ZEKANIN KULLANILMASI

Oğuz Ümit TAMER

Bruce D. MCDONALD III

Farouk HEMICI

Georgia KONTOGEORGA

GENİŞLETİLMİŞ ÖZET

Yapay Zekânın Kamu Sektörü Risk Yönetimindeki Uygulamaları

Günümüzde kamu idareleri, makine öğrenimi (ML), doğal dil işleme ve öngörücü analitik gibi yapay zekâ teknolojilerini, refah yardımları, vergi uyumu ve dolandırıcılık tespiti gibi alanlarda kullanmaktadır. Bu uygulamalar, büyük veri akışlarını işleyerek anomalileri hızlıca yakalama ve kaynak tahminleri yapma imkânı sunar. Örneğin, denetim fonksiyonlarında Yapay Zekâ Destekli Sürekli Denetim (Continuous Auditing) modelleri sayesinde tüm işlem verileri gerçek zamanlı izlenebilmekte ve denetçiler riskli işlemlere odaklanabilmektedir. Benzer şekilde, ML temelli tahmin teknikleri salgınlar veya sosyal hizmet talebi gibi ihtimalleri önceden öngörerek personel ve malzeme önlemleri alınmasına yardımcı olur. Bu sayede kamu operasyonlarında hizmet aksama riskleri azalmakta ve verimlilik artmaktadır. Yine denetim süreçlerinde derin öğrenme ve kümeleme algoritmaları, insan denetçilerin gözden kaçırabileceği karmaşık anomalileri tespit ederek denetim kalitesini yükseltir.

Temel Zorluklar: Algoritmik Şeffaflığın Olmaması, Önyargı ve Hesap Verebilirlik Eksiklikleri

Yapay zekâ (YZ) uygulamalarının avantajlarının yanı sıra ciddi riskleri de vardır. Derin öğrenme modelleri gibi bazı YZ sistemleri, karar verme süreçleri insanların anlayamayacağı şekilde gizli (black-box) çalışır. Bu şekilde algoritmik şeffaflık yoksunluğu, denetçilerin nasıl sonuç elde edildiğini izlemelerini güçleştirir ve hata riskini artırır. Ayrıca, geçmişe dayalı eğitilmiş modeller toplumsal önyargıları öğrenip pekiştirerek azınlık gruplarına karşı ayrımcılığa neden olabilir. Örneğin bazı Avrupa ülkelerindeki refah yolsuzluğu tespit sistemleri, azınlıkları orantısız biçimde hedef olarak adaletsizlik yaratmıştır. Son olarak, bu sistemlerde hesap verebilirlik boşlukları oluşmaktadır. YZ uygulamalarının neden yanlış sonuç verdiği veya hatalı kararlar aldığı anlaşılamadığı zaman, sorumluluk belirsizleşir. Bu sorunlar, yapay zekânın kamu sektöründe kullanılmasında şeffaflık, adalet ve insan denetimi gibi yönetim ilkelerinin eksik uygulandığının göstergesidir.

Düzenleyici Çerçevesler: AB Yapay Zekâ Yasası ve NIST AI RMF

Bu riskleri gidermek için uluslararası düzenlemeler geliştirilmiştir. Avrupa Birliği, yüksek riskli kamu uygulamalarına yönelik risk temelli bir düzenleme olan AB Yapay Zekâ Yasası'nı yürürlüğe koymuştur. Yasa, özellikle kamu programlarında kullanılacak AI sistemlerine şeffaflık, adalet ve insan denetimi gibi gereklilikler getirir. Benzer biçimde, ABD Ulusal Standartlar ve Teknoloji Enstitüsü (NIST) "AI Risk Management Framework" adlı bir kılavuz yayımlamıştır. NIST çerçevesi, Yönetim (Govern), Ölçüm (Measure) ve İdare (Manage) gibi işlevlerle YZ riskini tanımlama ve izleme süreçlerine odaklanır. Ayrıca OECD'nin Yapay Zekâ İlkeleri ve G7'nin kamu sektörüne yönelik YZ araç takımı gibi yönergeler, hükümetlerin sorumlu YZ stratejileri geliştirmesine destek olur. Bu düzenleyici girişimler, yapay zekâ tasarımına baştan risk yönetimi ve denetlenebilirlik mekanizmalarını süreçlerin ayrılmaz bir parçası haline getirme gerekliliğini vurgular.

Algoritmik Denetimin Rolü

Algoritmik denetim, yapay zekâ sistemlerini önceden belirlenmiş etik ve teknik kriterlere göre inceleyen bir yöntemdir. Uygulamayı geliştiren kurumun iç denetimleri ve bağımsız kuruluşların denetimleri şeklinde iki türü vardır. Tasarım ve eğitim aşamasında veri kalitesi, temsiliyet ve önyargı odaklı denetimler yapılır. Geliştirme aşamasında, model kartları ve veri seti föyleri gibi dokümantasyon araçları kullanılarak amaç, performans ve adil kullanım bilgilerinin kayıt altına alınması sağlanır. Yaygınlaştırma ve izleme aşamasında ise YZ sistemi sürekli takip edilerek istenmeyen etkiler ve ayrımcı sonuçlar hızlıca tespit edilir. Algoritmik denetim mekanizmaları sayesinde, kamu sektörü YZ sistemlerinin toplum değerlerine uyumlu kalması sağlanır ve risk algılama çerçevesi etkin bir şekilde hesap verebilirlik paradigmasına dönüştürülür. Bu süreç, tespit edilen risklerin yönetim yapıları içinde geri besleme olarak kullanılmasına olanak tanır.

Denetim Kurumları ve Kamu Yönetişimine Etkileri

Çalışmanın bulguları, özellikle politika yapıcılar, yüksek denetim kurumları ve iç denetim birimleri için yol göstericidir. Öncelikle, yapay zekâ araçlarının güvenli kullanılabilmesi için yönetim kurallarının uygulamaya alınmadan önce hayata geçirilmesi gerektiği vurgulanır. Yani YZ sistemlerinin denetlenmesi yalnızca sonuçların incelenmesiyle sınırlı kalmamalı, aynı

zamanda bu sistemlerin altyapısını ve işleyiş süreçlerini yöneten politikalar da değerlendirilmelidir. İkinci olarak, denetim faaliyeti sürekli olmalıdır; geleneksel, periyodik denetimler ML modellerinin dinamik yapısını yakalayamaz. Bu nedenle yüksek riskli kamu AI uygulamalarında Sürekli Yapay Zekâ Denetimi (Continuous Auditing of AI) standart hale gelmelidir. Üçüncü olarak, yapay zekâ insan işgücünü tamamen ortadan kaldırmayacak; yüksek hacimli veri analiz ve anomali belirleme görevlerini AI yaparken, insan denetçi ise etik değerlendirmeler, adalet kararları ve vatandaş etkileşimi gibi stratejik rollere odaklanacaktır.

Bu dönüşüm, denetçilerin dijital ve etik yetkinliklerini artırma ihtiyacını doğuracaktır. Uzun vadede, kamu kurumları yapay zekâyı bir ihtiyaçtan ziyade sorumlu bir gözetim süreciyle rutin şekilde kullanmalı; böylece YZ'nın ortak refaha katkı potansiyeli azami seviyede değerlendirilebilecektir. Türkiye'de Sayıştay gibi denetim kurumlarında da yapay zekâ araçları henüz başlangıç aşamasındadır, ancak potansiyel faydaların sınırsız olduğu düşünülmektedir. Uyum sağlamak için ulusal ve uluslararası rehberler dikkatle izlenmekte ve yukarıdaki risklerin önlenmesine yönelik tedbirler sıkı bir şekilde uygulanmaktadır.