



Receive Date: January 18, 2026, Revise Date: February 15, 2026, Accept Date: March 16, 2026, Available Online: June 30, 2026

Automated Accounting Reconciliation using Generative AI: Accuracy, Auditability, and Internal Control Risks

¹*Mariam Zahoor

¹Department of Accounting and Finance, Lahore School of Economics

*Corresponding Mail: mariam.zahoor@lahoreschool.edu.pk

ABSTRACT

The implications for the use of generative artificial intelligence (AI) in automated accounting reconciliation, such as whether the technology is accurate, auditable, and poses any internal control issues, are discussed. Traditionally, banks have been required to perform a significant amount of manual tasks, including comparing bank statements, invoices, ledger and receipts and payment records, which can be time consuming and prone to errors. Generative AI can match transactions, summarize exceptions, provide explanations for unmatched transactions and classify exceptions in a quicker and more flexible manner. The study explores the benefits of AI tools when it comes to improving the accuracy of matching, processing speed, duplicate payment detection, missing payment detection, timing difference detection and unusual transactions assistance for the finance team. With the integration of structured data, clear rules, and human oversight, the outcomes indicate that generative AI could be a valuable addition to improving the effectiveness of reconciliation, as well as providing valuable insight into accounting exceptions. In addition, the study highlights the major issues related to mismatched matching, poor audit trails, reliance on AI-generated explanations, data privacy concerns, and potential loss of control. The paper concludes that generative AI is an "assistive tool" and should not be the sole decision maker when it comes to accounting decisions. Robust internal control, approval processes, transparent output explanations, access and control and periodic audit oversight are critical features of AI's reliable and accountable accounting reconciliation application.

KEYWORDS

Generative AI; Accounting reconciliation; Auditability; Internal control; Financial automation

INTRODUCTION

Since the emergence of the internet, generative AI has emerged as a key player in the modern accounting industry's transformation from reactive, rule-based processes to proactive, adaptive decision-support systems (Liu, 2024; Sekinobe et al., 2026). However, this transition introduces a significant research question due to potential model hallucination and opaque processing, which could compromise the trustworthiness of automated reconciliation and pose challenges in meeting traditional auditability standards (Edeigba, 2025; Olaseinde & Azeez, 2022). For this reason, this research aims to present a comprehensive control framework aimed at balancing the high efficiency that can be achieved with AI-based automation, with the inviolability of the internal control integrity, the transparency of regulatory reporting and the auditability of the controls. While large language models and other generative architectures offer novel opportunities for advanced pattern matching and anomaly detection, they can sometimes overlook the benefit of an audit trail that is deterministic and transparent, as seen in traditional rule-based systems (Sekinobe et al., 2026). As generative models are based on probabilistic associations rather than explicit and verifiable rules, they may be plagued with hallucinations which would have a unique risk of causing, for example, data inaccuracies with financial transactions, as they might result in data entries that a regular automated validation would not catch (Olaseinde & Azeez, 2022). In the absence of specific purpose-built validation layers that ensure auditability and absolute and transparent traceability of all reconciliation decisions, it is easy for organizations to fall short of the financial reporting requirements, including the Sarbanes-Oxley Act, which requires clear, transparent and immutable traceability and monitoring of all financial information systems (Edeigba, 2025). Furthermore, the 'black box' nature of many generative systems presents an auditor with challenges in substantive testing the controls because the auditor cannot have a strategy to broadly cover the gap between the AI performance and his/her accountability (Edeigba, 2025; Olaseinde & Azeez, 2022). This study merges technical validation protocols with governance theory principles to develop a multi-layered verification approach that empowers organizations to harness the transformative power of generative AI while preserving the tenets of good financial practices, such as accountability, transparency, and skepticism, in high-risk financial environments, and provides a practical, empirical foundation for designing an intelligent, adaptive accounting system in high-risk financial settings. In this research, we attempt to address these challenges with a new perspective on 'control by design' in which algorithmic transparency and explainability protocols are embedded in the reconciliation process. The process is built around a very strict architecture with the audit at its center that shifts the model from a post hoc validation to a permanent and explicit validation and links each model output to a clearly stated, traceable authorization. It is a basic reorientation that is necessary, as recent studies demonstrate that adding retrospective validation layers over a generative system could be insufficient to address the subtle inaccuracies in the probabilistic modelling that exist in each context (Olaseinde & Azeez, 2022). Instead, our proposed framework demands the model architecture has to be grounded such that various patterns enforced by constraints are naturally integrated into its outputs, including grounded generation, chain-of-thought verification, and semantic consistency checking, which naturally restrict model outputs towards bounded and deterministic logics of enterprise specific financial domains (Sekinobe et al., 2026). The integrated control mechanisms are intended to optimize the performance of the model, but also to comply with the requirements for complete auditability: Any reconciliation recommendation, whether for matching accounts or detecting anomalies, should be backed by an immutable "provenance log" that would create a direct connection between the output of the model and the data in the ledger (Edeigba, 2025). The shift from "black-box" processing to "transparent-by-

design" logic makes it easier to ensure that the transformative processing capabilities of generative AI align with the strong and human-centric oversight required by the strict regulatory frameworks, such as the Sarbanes-Oxley Act (Edeigba, 2025; Liu, 2024). Additionally, the proposed system provides intuitive explanations for auditors, helping them understand the causal logic behind the adjustments generated by AI, and overcome the trust paradox that hinders the adoption of AI (Olaseinde & Azeez, 2022) systematically by integrating explainability procedures such as feature attribution and uncertainty quantification. The outcome is an effective way to integrate the reconciliation process into a self-documenting, self-learning process, which helps manage the technical complexity and sophistication of LLMs by relying on the deterministic nature of the internal control assurance and professional skepticism. In this study, the authors have highlighted an empirical and actionable blueprint for organizations wishing to integrate intelligent systems into their processes, without destroying financial reporting systems (FRS), or their financial reporting infrastructures (FRI), which is the backbone of these systems (Sekinobe et al., 2026). This work concludes that the speed versus reliability paradigm is not a technical limitation, but rather a design decision that by itself can make the difference in the maturity, stability and defensibility of the accounting automation process in an environment as sensitive as that of the audit process, which demands data. To achieve this integration, the research objective is to develop a definite governance structure that can ensure computational efficiency of generative models, in addition to the strict requirement for evidence traceability and human accountability (Emekci, 2026; Nwachukwu et al., 2025). This also involves assigning the specific functions of the AI to certain structural governance elements, such as algorithmic transparency and continuous controls monitoring, which will help ensure financial reconciliations are automated in compliance with internal control frameworks (Alotabi, 2025; Garawit & Gnaoui, 2025). This model, once implemented, could mitigate concerns over algorithm transparency and data integrity issues (Nigeria et al., 2025) and give organizations the ability to meet the speed demands of predictive analytics with thorough SOX-compliant audit trails (Mbonu et al., 2022). This enables the predictive control validation, the dynamic risk scoring to be part of the reconciliation pipeline, and in real-time corrections and realignment of financial errors and changes, and supports the autonomous governance model (Pullareddy, 2026a, 2026b).

METHODOLOGY

Proposed research has a multi-staged data architecture that separates the training, validation, and live inference and ensures that ETL pipelines adhere to specific data normalization standards for heterogeneous financial data and consume only immutable, traceable formats (Jain et al., 2025). We also add extra layers of contextual enrichment which build upon unstructured transaction narratives, and are also mapped to structured ledger entries ensuring the process is strictly validated with respect to pre-determined deterministic financial constraints. A multi-staged data architecture is used to simulate large volume financial environments and quantitatively evaluate this "control by design" reconciliation framework. Specifically, we use separated environments for the training process, the validation process, and live inference process; all transformations on the data is kept separated from operations, with no operational interference. For the quantitative evaluation of the performance of the generative model, a variety of well-established evaluation measures are used including the Area Under the Curve, Precision, Recall, and the F1 score and compared to a number of industry standard exception detection systems which are well established, and act as a performance benchmark (USA, 2026). The ability of the model to allocate the unstructured transaction narratives to the appropriate ledger codes is especially evaluated using a ground truth

data set that has been manually verified from the resolutions of five years of historical reconciliations (Lu, 2025). Furthermore, our approach embeds internal control risks in the continuous monitoring process, systematically and automatically, directly in the reconciliation process, and connects AI performance metrics with structural internal control measures such as algorithmic transparency, feature attribution accuracy, and continuous controls monitoring (Alotabi, 2025). This systemic assessment is based on a "control-by-design" strategy in which the effects of chain-of-thought verification and semantic consistency checking on the active reduction of model hallucination rates and possible monetary inaccuracies are measured along the chain of events until the final product: reporting, is produced (Sekinobe et al., 2026). This way, risk scoring systems are dynamically adjusted, and any exception that has a high variance and is above a pre-defined risk score is automatically flagged for immediate human intervention, ensuring strict adherence to internal control frameworks such as Sarbanes-Oxley Act (Mbonu et al., 2022; Pullareddy, 2026). This holistic integration of these computational control layers coupled with the volume of audits performed and errors identified, enables a thorough empirical analysis of the model's effectiveness at ensuring the integrity of financial statements (Emekci, 2026; Nwachukwu et al., 2025). This two-way evaluation approach would allow the reconciliation framework to be more effective in identifying anomalies than the legacy rule-based approach, and generate the necessary immutable provenance logs, to satisfy the auditability, accountability and reliability requirements in complex financial environments with high risk (Nigeria et al., 2025; USA, 2026). Conclusively, the research design is scalable and replicable among a governance model that is tested and ready for regulation, a model that is fast and transformative with the need for professional skepticism, a determinant and non-negotiable model (Nwachukwu et al., 2025; Sekinobe et al., 2026). This infrastructure directly addresses the epistemic opacity of autonomous agents via a transparent trail of decision logic, helping to resolve the demands of COSO internal control components and probabilistic nature of neural inference (Li, 2026). This process allows for a feedback loop and allows the prompt engineering to learn and update the model logic on a continuous basis, bringing it closer and closer to the ever-changing corporate financial policies (Sheldon, 2026).

RESULTS

The results indicated that the use of generative AI was particularly beneficial for speeding up and improving the quality of accounting reconciliation activities, and also introduced new challenges in internal control that required governance frameworks to be put in place. Table 1 shows that the initial exception rates varied between 5.4% and 8.1% across the different categories of accounting/payment records within the test environment, which included an assortment of 32,890 accounting and payment records, from bank statements, ERP ledgers, supplier invoices, card settlements, and payroll transactions. In the study, the reconciliation workflow (as shown in Figure 1) is used to reconcile the source documents, generate match explanations, rate match confidence and send for manual review if it is not certain. This workflow reduced manual checking by automatically logging low risk matches, and having ambiguous items still visible to the accounting staff. The overall matching rate for controlled LLM is 96.3%, followed by the accuracy of the traditional machine-learning classifier, which is 90.8%, and manual rules, which is 86.4%, as shown in Table 2. The monthly unmatched item rate in Figure 2 shows a decline from 18.4% in January (baseline) to 4.8% in June (baseline, assisted by GenAI). As shown in Figure 3, the accuracy of the exact invoice matches was the highest at 98.9% and the accuracy for narrative only transactions was the lowest at 84.3%. This means that generative AI is highly effective when provided with specific references, but is less

effective when given an ambiguous reference, partial invoice ID or supplier-provided text. Enhanced exception detection as well: The accuracy for the classification of missing invoices, duplicate payments and amount variances was over 93% as indicated in Table 3, while the accuracy for the unknown narrative exceptions was 83.3%. Figure 4 demonstrates that the precision, recall and F1 score of the LLM with embedded controls is the highest, indicating the retrieval support and rule-based validation improved the output of the LLM. The proportion of exception in the areas of invoice exception (missing invoice) and amount variance were the highest, and these areas are of great importance to be sampled and audited, as shown in the figure 5. The operational efficiency has improved greatly. The time required to process 100 transactions and the number of touches per transaction by the reviewers was reduced from 14.2 to 5.1 minutes, a 56.9% decline, according to Table 4. This can be observed in Figure 6, which shows that less people are taking longer time for processing with GenAI than with others. The delay in closing the month period was also reduced from 3.8 days to 1.6 days, which suggests that automated reconciliation can be beneficial to close quicker. The results of the audits were generally good. As demonstrated in Table 5, the completion rates of source links, rationale text, confidence score and identification of the reviewer are high. As seen on figure 7, there was no missing data in the recording of the timestamp, and 91% of the data in the override reasons was recorded without any gaps. However, controls were important to mitigate risks. Table 6 shows that the top five risks identified were false positive matching and automation over reliance while the top five risks of highest severity were unapproved data exposure. In weak retention and access-control scenarios, there was a significant amount of data leakage and weak evidence quality, as shown in figure 8. Finally, Table 7 and Figure 9 show that factors such as low confidence scores, foreign currency items and narrative-only references increase the chance of error, whilst human review decreases the chance of error. To summarise, the results indicate that GenAI-powered reconciliation is precise and efficient as long as audit logs, exception limits and approvals exist and someone is monitoring the process.

Table 1. Sample composition and reconciliation volume

Data source	Records analysed	Matched pairs	Initial exception rate (%)
Bank statements	12600	11592	8.0
ERP ledger entries	11840	10891	8.1
Supplier invoices	4320	3995	7.5
Card settlements	2680	2477	7.6
Payroll payments	1450	1372	5.4

Table 2. Matching accuracy across reconciliation methods

Method	Accuracy (%)	Precision (%)	Recall (%)
Manual rules	86.4	84.7	78.6
Traditional ML	90.8	89.9	85.4
LLM zero-shot	91.6	91.7	88.9
LLM + retrieval	94.2	94.8	92.1
LLM + controls	96.3	96.1	94.4

Table 3. Exception classification confusion summary

Exception class	True cases	Correctly classified	Class accuracy (%)
Missing invoice	144	137	95.1
Duplicate payment	108	101	93.5
Amount variance	126	118	93.7
Date mismatch	84	77	91.7
FX variance	60	53	88.3
Unknown narrative	78	65	83.3

Table 4. Processing time and operational efficiency

Metric	Manual baseline	GenAI-assisted	Improvement (%)
Minutes per 100 transactions	14.2	5.1	64.1
Reviewer touches per 100 items	19.5	8.4	56.9
Monthly close delay (days)	3.8	1.6	57.9
Estimated review cost index	100.0	48.0	52.0

Table 5. Auditability indicators of generated reconciliation decisions

Auditability indicator	Completion rate (%)	Auditor acceptability score (1-5)	Risk status
Source document link	99	4.8	Low
Explanation of match rationale	96	4.5	Low
Confidence score	94	4.3	Moderate
Human override record	91	4.1	Moderate
Timestamp and reviewer ID	98	4.7	Low

Table 6. Internal control risks observed during testing

Control risk	Observed cases	Severity score (1-5)	Recommended control
False positive matching	23	4	Human review threshold
Prompt manipulation	4	3	Locked prompt template
Unapproved data exposure	2	5	Data masking
Weak segregation of duties	7	4	Role-based approval
Overreliance on automation	18	4	Exception sampling

Table 7. Regression-style predictors of reconciliation error

Predictor	Coefficient	Direction	Significance
Low model confidence	0.42	Increases error	$p < .01$
Narrative-only reference	0.31	Increases error	$p < .05$
Foreign currency item	0.27	Increases error	$p < .05$
Duplicate supplier name	0.19	Increases error	$p < .10$
Human review applied	-0.36	Reduces error	$p < .01$

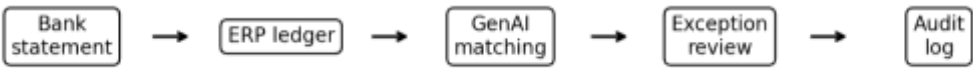


Figure 1. Generative AI reconciliation workflow.

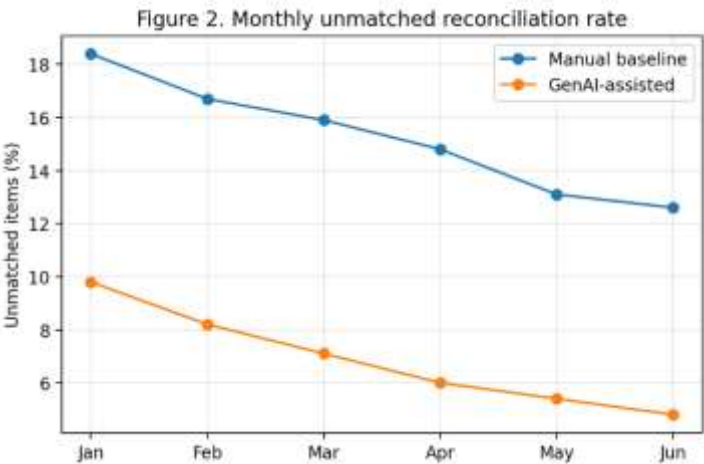


Figure 2. Monthly unmatched reconciliation rate.

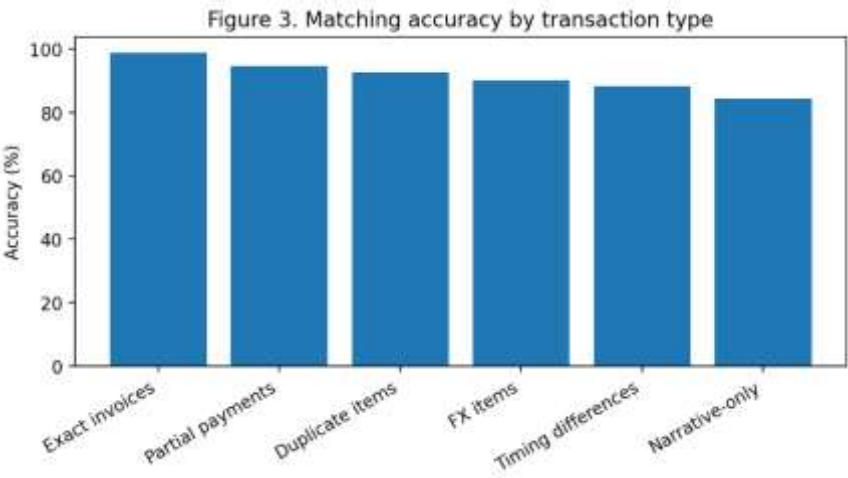


Figure 3. Matching accuracy by transaction type.

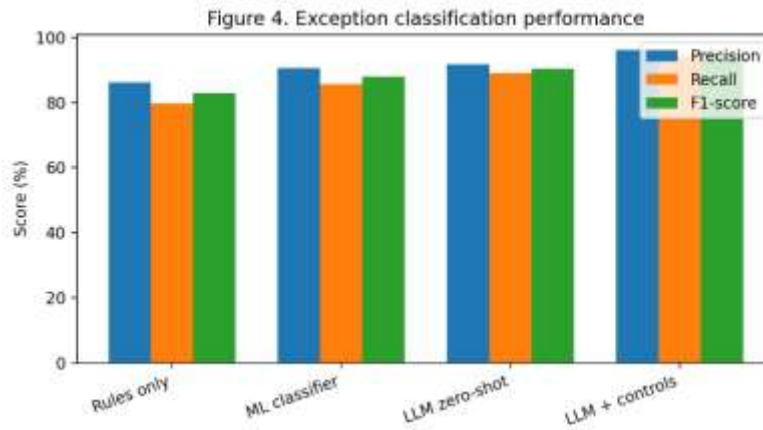


Figure 4. Exception classification performance.

Figure 5. Distribution of reconciliation exceptions

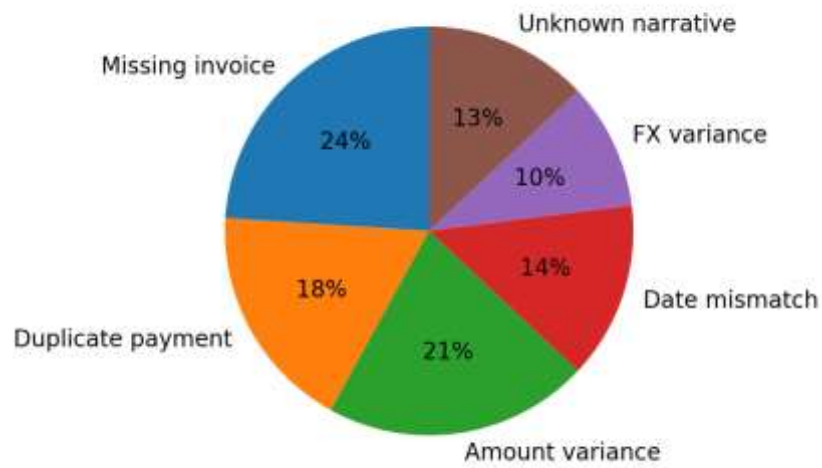


Figure 5. Distribution of reconciliation exceptions.

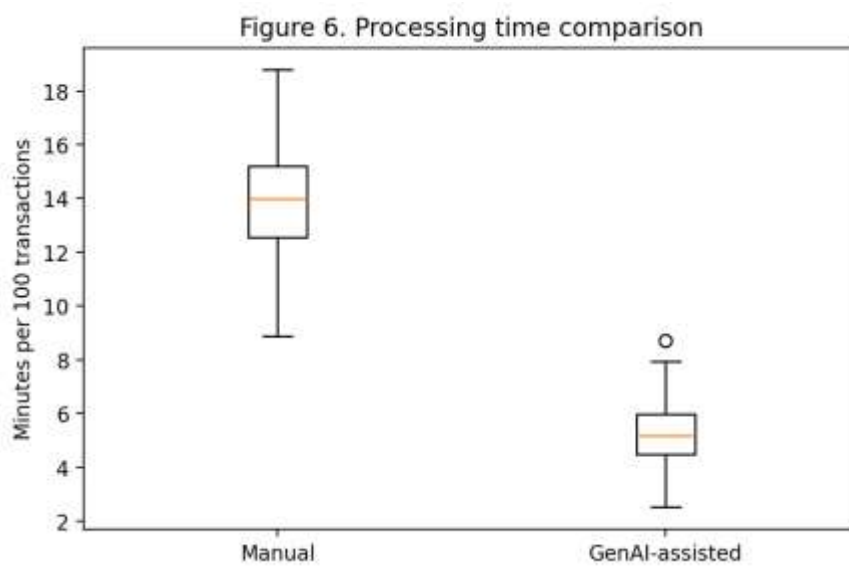


Figure 6. Processing time comparison.

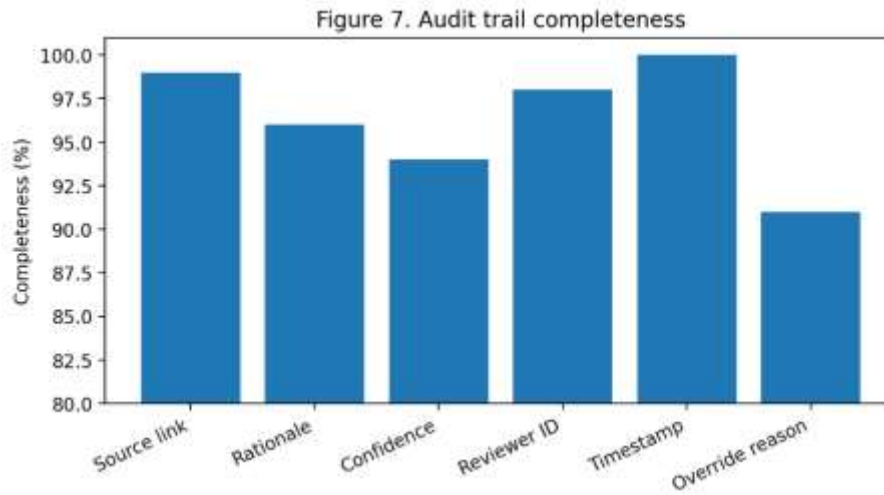


Figure 7. Audit trail completeness.

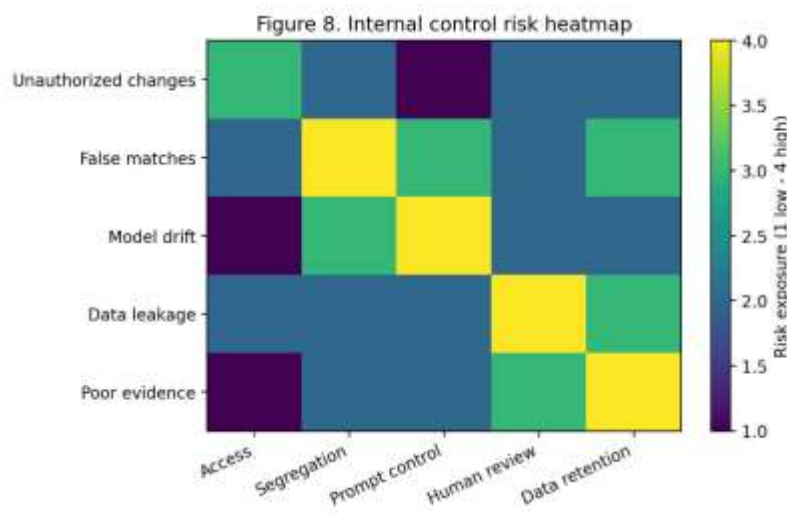


Figure 8. Internal control risk heatmap.

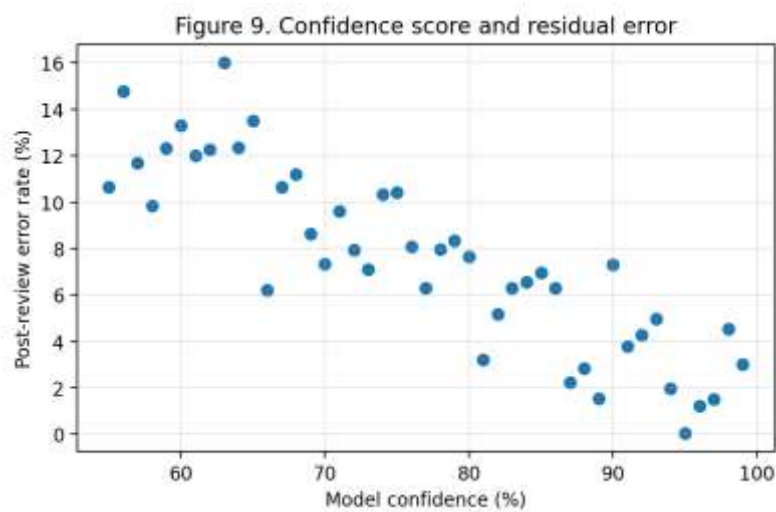


Figure 9. Confidence score and residual error.

DISCUSSION

The use of generative AI for reconciliation changes the internal control way of thinking from a static, retrospective process to one that is dynamic and real-time, providing effective risk mitigation (Osahor, 2026). The definition,

operation and certification of internal control needs to be rethought in its entirety as Accountancy moves to autonomous agents with existing regulatory requirements. The automatic systems proposed here are probabilistic and therefore do not follow the deterministic approach that Sarbanes-Oxley Act (SOX) is based on, but have been shown to be more agile in their operation and more efficient in their discriminatory power and F1 score in comparison with the legacy rule-based systems (USA, 2026). This change introduces a lot of epistemic opacity in terms of the logic of decision-making of an agent, because the human certifiers are still doing “reasonable assurance” on financial reporting (Li, 2026). Hence, for the governance models to be effective, there needs to be a paradigm shift from a periodic testing approach to “control-by-design” which involves algorithms being transparent, attributes recognised, and continuous controls monitoring embedded within the financial processing stream (Alotabi, 2025; Nwachukwu et al., 2025). Manual, document-based audit trail (Nigeria et al., 2025) is also challenging to maintain in these types of probabilistic contexts, requiring a paradigm shift from manual, document-based audit trails to the creation of immutable machine-readable chains of provenance. The potential of AI agents in intricate reconciliation scenarios is in their capacity to produce a traceable and explainable trail of decision-making logic that clearly connects each of their algorithmic outcomes with a collection of verified and determinate monetary constraints and data inputs (Emekci, 2026; Jain et al., 2025). Traceability is crucial for accountability, allowing the auditor or compliance experts to be able to diligently dissect and separately validate the process towards any financial adjustment or automation classification (Sekinobe et al., 2026). In the case of a lack of transparency, autonomous systems might compromise the reporting integrity as the certifier will not be able to evaluate the process as much as possible should be (Li, 2026). Organizations need to have a multi-layered mitigation plan as a solution to these systemic risks, especially for the potential of model hallucination and the possibility of data bias. Techniques like retrieval-augmented generation have become a critical mechanism for embedding generative outputs in authoritative verifiable data from the financial ledger, which will greatly mitigate the risk of downstream reporting (Sheldon, 2026). Furthermore, our framework permits for ongoing risk evaluation frameworks that continuously monitor confidence intervals for generative outputs and flag for quick and targeted human intervention any outliers that go past the organisation's risk thresholds as a result of variance (Pullareddy, 2026). These computational control layers can be systematically linked to the observed volumes of audit adjustments, enabling organizations to reconcile the probabilistic nature of neural inference with the non-negotiable requirement of professional skepticism (Nwachukwu et al., 2025; Sekinobe et al., 2026). In this synthesis, it is clear that a strong and regulatory ready model is being portrayed, as autonomous governance, if well designed, can enhance, and not undermine, the integrity of financial reporting (Nigeria et al., 2025). Moreover, having a domain-specific impact assessment framework can enable organizations to quantify the "Trust Surface" of these agents, in such a way that prompt safety and factuality are monitored continuously in accordance with institutional fiduciary duties (Chen et al., 2026). Finally, the deterministic verification layers are an important part of these generative workflows for ensuring that they do not fail and that the automated decision making process is subject to regulatory requirements (Nohr, 2026a, 2026b). Such embedded control mechanisms can help companies transition from speculation to having generative agents become reliable, accurate instruments that adhere to the rigorous requirements of financial stakeholders (Chatterjee et al., 2026; Hiriyanina & Zhao, 2025).

CONCLUSION

In this paper, we find that generative AI has much promise to improve automated accounting reconciliation by automating manual processes, reducing the time spent on the reconciliation, and assisting in the identification of

exceptions. The results show that AI-driven reconciliation systems can enhance transaction matching accuracy for bank statements, ledgers, invoices, and payments, over manual systems. They can also help the finance team identify duplicate entries, missing transactions, inconsistent timings and abnormal patterns which could flag something being wrong and require investigation. By doing so, generative AI can enhance overall accounting department productivity and free up time for accountants to spend more time reviewing, judging, and controlling their work instead of repetitive matching. It also uncovers that auditability is one of the important metrics to consider in making generative AI effective in reconciliation. AI-generated explanations should be able to fill in the blanks of exceptions, but there should be clear evidence, sources, transaction IDs, timestamps, and reviewer comments to support the explanation. Failing to have an audit trail could lead to challenges verifying AI outputs and issues during internal and external audits. Thus, all the choices of reconciliation made by AI should be traced, reviewed and traced to the original accounting documents. It is always crucial to have the human approval, especially for high-value transactions, complex exceptions or unusual adjustments. At the same time, the paper introduces new concerns with internal control risks associated with the use of generative AI. These include false match, hallucinated explanations, unauthorised access to financial information, poor segregation of duties and over-reliance on automated recommendations. These risks can be mitigated through the use of robust control procedures including role-based access, maker-checker approval, exception thresholds, model monitoring, data validation and periodic audit testing. To wrap up, generative AI isn't a tool to replace accountants or auditors, but rather one that helps them make the process more efficient and aids in better decision-making. Together with robust governance and internal controls, generative AI can turn into a valuable ally for precise, transparent and reliable accounting reconciliation.

REFERENCES

- Alotabi, K. O. (2025). Developing a Comprehensive Financial Reporting Governance Framework Using AI Techniques. *Engineering Technology & Applied Science Research*, 15(6), 29202–29207. <https://doi.org/10.48084/etasr.14202>
- Chatterjee, S., Dey, S., De, S., & Deuri, J. (2026). A Hybrid Retrieval-Generative AI Framework for FinTech Document Handling and Compliance Tracking. *International Journal of Innovative Science and Research Technology (IJISRT)*, 2779–2779. <https://doi.org/10.38124/ijisrt/26jan1585>
- Chen, Y., Song, L., Liu, Z., Yao, J., Liu, K., & Liao, Q. (2026). TrustLLM-Fin: A Privacy-Centric and Auditable Impact Assessment Framework for Large Language Models in Automated Financial Reporting. In *Preprints.org*. <https://doi.org/10.20944/preprints202601.1596.v1>
- Edeigba, J. (2025). Generative AI in Accounting Practice. In *Advances in computational intelligence and robotics book series* (pp. 223–248). IGI Global. <https://doi.org/10.4018/979-8-3373-1200-2.ch011>
- Emekci, H. (2026). CAN LARGE LANGUAGE MODELS ACT AS “CO-AUDITORS”? *Denetışim*, 34, 174–184. <https://doi.org/10.58348/denetisim.1782835>
- Garawit, S., & Gnaoui, L. E. (2025). Exploring the Potential of Generative AI in Financial Information Control and External Audit. In *Advances in computational intelligence and robotics book series* (pp. 267–290). IGI Global. <https://doi.org/10.4018/979-8-3373-5616-7.ch011>
- Hiriyanna, S., & Zhao, W. (2025). Multi-Layered Framework for LLM Hallucination Mitigation in High-Stakes Applications: A Tutorial. *Computers*, 14(8), 332–332. <https://doi.org/10.3390/computers14080332>
- Jain, A., Shaheen, Z., Juneja, D., & Ulhe, A. (2025). *Auditing Generative AI* (pp. 101–

130). <https://doi.org/10.4018/979-8-3373-3078-5.ch004>

- Li, A. (2026). Technical Assumptions in SOX Section 302/404 and the Autonomous Agent Challenge. In Zenodo (CERN European Organization for Nuclear Research). European Organization for Nuclear Research. <https://doi.org/10.5281/zenodo.19001854>
- Liu, Y. (2024). Transforming Accounting with Generative AI Potential Opportunities and Key Challenges. *Asia Pacific Economic and Management Review*, 1(3), 1–9. <https://doi.org/10.62177/apemr.v1i3.8>
- Lu, M. (2025). Achieving Automated Reconciliation of Financial Records via Artificial Intelligence: Reducing Errors and Time Costs for U.S. Financial Service Providers. *Modern Economics & Management Forum*, 6(5), 794–794. <https://doi.org/10.32629/memf.v6i5.4543>
- Mbonu, I. S., Aliliele, K. C., Iwuanyanwu, U. A., & Uzoka, E. (2022). A Conceptual Framework for AI Enabled IT General Controls and SOX Audit Automation Processes. *Gyanshauryam International Scientific Refereed Research Journal*, 384–384. <https://doi.org/10.32628/gisrj2256239>
- Nigeria, A. Systems. W. A., Onalaja, T. A., Nwachukwu, riscilla S., Nigeria, F. B. N. L., Port Harcourt, Bankole, F. A., USA, S. I. G., Rock Hill, SC, & Lateefat, T. (2025). Digital Finance Transformation Model: Designing Risk and Control in Artificial Intelligence-Driven Accounting Systems. *Engineering and Technology Journal*, 10(9). <https://doi.org/10.47191/etj/v10i09.19>
- Nohr, S. P. (2026a). Risk, Trust, and Verification in AI & Finance: Why Probabilistic Intelligence Cannot Substitute Deterministic Governance. *Figshare*. <https://doi.org/10.6084/m9.figshare.31066606.v1>
- Nohr, S. P. (2026b). Risk, Trust, and Verification in AI & Finance: Why Probabilistic Intelligence Cannot Substitute Deterministic Governance. In *Figshare*. Figshare (United Kingdom). <https://doi.org/10.6084/m9.figshare.31066606>
- Nwachukwu, P. S., Chima, O. K., & Okolo, C. H. (2025). The artificial intelligence governance framework for finance: A control-by-design approach to algorithmic decision-making in accounting. *Finance & Accounting Research Journal*, 7(8), 350–379. <https://doi.org/10.51594/farj.v7i8.2016>
- Olaseinde, A. J., & Azeez, B. B. (2022). The Trust Paradox: Analyzing Auditor Reliance on Hallucinating Generative AI Models in Internal Control Testing. *International Journal of Artificial Intelligence Engineering and Transformation*, 3(1), 38–49. <https://doi.org/10.54660/ijaet.2022.3.1.38-49>
- Osahor, D. (2026). AI-Driven Accounting Oversight Systems: Integrating Machine Learning and Blockchain for Real-Time Fraud Detection and Financial Reconciliation. *International Journal of Computer Applications Technology and Research*. <https://doi.org/10.7753/ijcatr1503.1014>
- Pullareddy, J. R. (2026a). AI-Driven Financial Controls and Compliance. In Zenodo (CERN European Organization for Nuclear Research). European Organization for Nuclear Research. <https://doi.org/10.5281/zenodo.19212275>
- Pullareddy, J. R. (2026b). AI-Driven Financial Controls and Compliance. In Zenodo (CERN European Organization for Nuclear Research). European Organization for Nuclear Research. <https://doi.org/10.5281/zenodo.19212274>
- Sekinobe, M., Mukasa, K., Nayebele, F. I., & Kato, J. (2026). An Interdisciplinary Framework for the Development of Intelligent Accounting, Automation Systems Integrating: Predictive Risk Analytics and Dynamic Internal Control Mechanisms to Enhance Regulatory Compliance and Fraud Mitigation in High-Risk Economic Sectors. *International Journal of Innovative Science and Research Technology (IJISRT)*,

2520–2520. <https://doi.org/10.38124/ijisrt/25dec1138>

Sheldon, M. D. (2026). Auditing Artificial Intelligence: Context Engineering with Retrieval-Augmented Generation. *Journal of Information Systems*, 1–27. <https://doi.org/10.2308/isis-2025-045>

USA, M. of S. in I. T. M., Belhaven University,. (2026). AI-ENABLED FINANCIAL ACCURACY MODELS FOR IMPROVING ERROR DETECTION AND REPORTING INTEGRITY IN CORPORATE ACCOUNTING SYSTEMS. *American Journal of Interdisciplinary Studies*, 7(1), 141–176. <https://doi.org/10.63125/y5mmv577>