

Maintaining Human Oversight Within Agentic Artificial Intelligence-Driven Workflows for Enterprise Operational Decision-Making and Automated Business Process Optimization Systems

Sai Viswa Teja Arumilli
Independent Researcher, USA.

Abstract: The human-in-the-loop agentic process of work in enterprise operational decision-making systems in this research. To enhance accuracy, reliability, and accountability in the applications of AI, the study will examine the advantages and disadvantages of automation, as compared to human supervision. The paper adopts a secondary data-driven research model to review the literature on the nature of agentic AI systems, how humans interface with these systems, and how enterprises make decisions. These results underscore the importance of having human control when incorporating it with AI-only solutions, as it boosts the quality of decisions, minimizes errors, and promotes accountability. The paper underscores that within the current environment of AI-embedded business operations and the digital age, hybrid decision-making approaches are important in the realization of ethical, scalable, and effective business operations in any enterprise.

Keywords: *Agentic AI, Human-in-the-loop, Enterprise decision-making, Artificial intelligence governance, Workflow automation, Risk mitigation, Ethical AI systems*

I. INTRODUCTION

Background

AI agentic workflow within industries is proving to be a radical change in the way enterprise operations are decided, bringing order to the operations in terms of both process simplification and automation of intricate jobs. However, in crucial business situations, the issues that are oriented towards accountability, transparency, and reliability are brought up due to the complete freedom of systems [1]. The combination of AWS, like with its services, Athena, and Redshift, and the integration of humans-in-the-loop conceptualization assists in keeping business decisions aligned with business goals, ethics, and regulations. Incorporating human control, AI-based operations can be enabled to make more holistic risk decisions and are more flexible to operate in fluid operational conditions [2]. The study seeks to investigate the issue of balancing and optimization of automation/ human intervention at the level of the enterprise to enhance more efficient system performance without compromising the trust, control and cautious utilization of intelligent systems in decision-making procedures at the enterprise model.

Research Aim

The aim is to examine human-in-the-loop agentic workflows, improving enterprise operational decision-making and governance effectiveness strategies models

Objectives

- To explore how agentic AI workflows integrate human oversight in enterprise operational decision-making processes

- To evaluate the effectiveness of human-in-the-loop systems in reducing operational risks and errors
- To identify issues and challenges in implementing agentic workflows within enterprise governance and compliance frameworks
- To recommend strategies for balancing automation and human judgment in enterprise decision-making systems effectively

Research Questions

1. How do agentic AI workflows integrate human oversight in enterprise operational decision-making processes effectively?
2. How effective are human-in-the-loop systems in reducing operational risks and errors significantly?
3. What issues and challenges exist in implementing agentic workflows within enterprise governance and compliance frameworks?
4. How can organizations balance automation and human judgment in enterprise decision-making systems effectively and sustainably?

Problem Statement

More and more frequently, AI is applied in an agentic style within enterprise processes to decide on operations, yet no frameworks are obvious to balance automation and human control [3]. This, of course, poses a number of accountabilities, transparency, principles of governance and ethics, and concerns the trust and the problems as a whole in an organization with complicated operational environment today.

Novel Contribution

In the current study, a systematic view is presented on one of the elements: human integration in the practice facet of agentic processes, with an intention to advance a debate on the benefits related to automation and monitoring of agentic processes [4]. It brings forward governance questions and proposes the improvement of decision support processes, thereby making the AI-infused systems more trusted, accountable and efficient.

Rationale

Companies are resorting to agentic types of AI systems to assist in their operations and the human supervision of AI is increasingly becoming necessary as it entails accountability. The significance of it is that the companies are now even more dependent on AI systems to make their decisions, yet require people to control them in order to be responsible. The ignorance about the advantages of the human-in-the-loop systems as far as governance and risk mitigation is concerned, is reduced by AI systems [5]. Today around the world the AI systems play a vital role in ensuring good usage of AI technology in complex organizational environments.

II. LITERATURE REVIEW

Human-in-the-loop integration in agentic AI enterprise decision workflows

Now more than ever before, human-in-the-loop has been identified as a key element to achieve robust enterprise decision-making by the agentic AI workflow. Human in the loop Systems with incorporation Human autonomous AI agents collaborate with a human to improve contextual understanding and accountability in sensemaking and execution [6]. Agentic Systems benefit in terms of consequences, rapidity and efficiency, however are unable to maintain an interpretable view in complex business environments.



Fig 1. Human–AI loop configurations

Outputs are validated and biases and alignment are corrected by human efforts against the objectives of the organization. The complete use of this integration in a business environment makes it possible to arrive at collaborative intelligence, which allows machines to work on repetitive tasks and humans to make strategic decisions [7]. Nonetheless, the issue of ensuring a

10.48047/jocaaa.2025.34.10.33

smooth coordination in applications remains in place, as it is challenging due to the nature of workflows and the application design. Defining roles, flexibly allowing interfaces, and promptly instating the feedback loops should be done to support trust in operational decision-making frameworks that involve AI, across industries.

Effectiveness of human oversight in reducing operational risks

Human oversight is a key component in minimizing operational risks in enterprise systems that are driven by AI. Monitors allow automated decision-making and detection of irregularities as well as prevention and compliance of organizational policies to be ensured [8]. Studies have indicated that man is a contributor of some of the best risk mitigation in high stakes areas such as finances, healthcare and logistics.

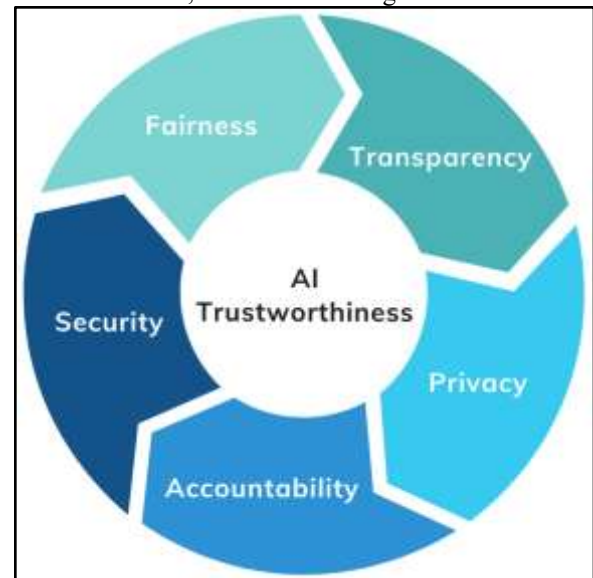


Fig 2. Components of AI trustworthiness

Unless agentic AI systems are managed correctly, they may yield biased or incorrect responses, increasing the likelihood of mistakes in working. Human validation of decisions and contextual reasoning are not possible with AI systems [9]. Excessive hand-picking is however haphazard and time-consuming. In this way, there will be a necessity in balanced supervision models in order to maximize risk abatement efficiencies and responsiveness. Good governance frameworks play an important role in providing reliability, accountability, and resilience to AI-supported enterprise processes.

Challenges in implementing agentic workflows within enterprise governance systems

A number of control, transparency and integration issues prevail regarding deploying agentic workflow in enterprise government. Organizations find it hard to incorporate an autonomous process of making decisions over AI into the existing regulatory and

compliance regimes [10]. When AI-generated decisions are not interpretable, it becomes challenging to audit and keep track of people's accountability.

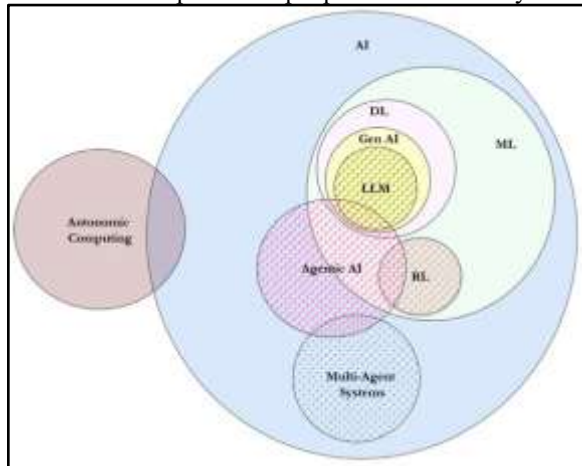


Fig 3. The Rise of Agentic AI

Additionally, changes in workflows to accommodate human control in automated systems are not straightforward and could be a big challenge. Not all employees are ready to embrace this and there is lack of skills [11]. Applications of AI agents are also associated with data privacy and the security of the information used. Moreover, one of the lifelong problems is the coherence of the decisions made by humans and machines [12]. These challenges require proper governance structures and controls on the existence of clear policies and continued supervision of the implementation of agentic workflows in the enterprise context.

Strategies for balancing automation and human judgment in enterprises

The tradeoff between automation and human judgment will be required in the enterprise context, which will then require the development of hybrid decision-making frameworks. These frameworks are based on the notion that the machine agents may advance to undertake simple tasks and data-intensive tasks, whereas human agents make the complex decisions [13]. The study insists on the need to have mechanisms in which human beings can step in whenever the content that AI hands out fall below expectations.

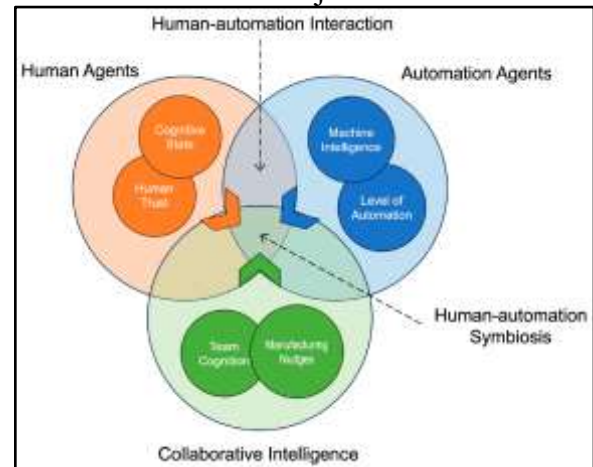


Fig 4. Human-Automation Interaction

More trust and the ability to conduct a better decision validation of transparent AI systems require transparency. Teaching that workers can do much with AI and are experiencing the limitations thereof is also beneficial when creating cooperation with AI [14]. Also, feedback mechanisms should have been provided, in order that the performance of the AI models would be constantly refined by the correction of humans. There should be clear roles, responsibilities and escalation procedures within governance policy [15]. Such a fair practice would allow a more efficient and effective decision-making process, reduce risk, and improve accountability throughout the AI-facilitated decision-making process at an enterprise, without losing sight of ethics and operation.

Literature Gap

Though much work has been done on either the AI automation or HIL systems there has been little work done on integrating both of these aspects as a single unit within the agency enterprise processes. Not only is the interaction of human autonomy on operational decisions in real-time not understood, it is also barely understood [16]. Also, a lack of good tests of governance models, risk-balancing mechanisms, and practices to effectively scale up to different types of enterprises and industries is evident.

10.48047/jocaaa.2025.34.10.33

III. METHODOLOGY

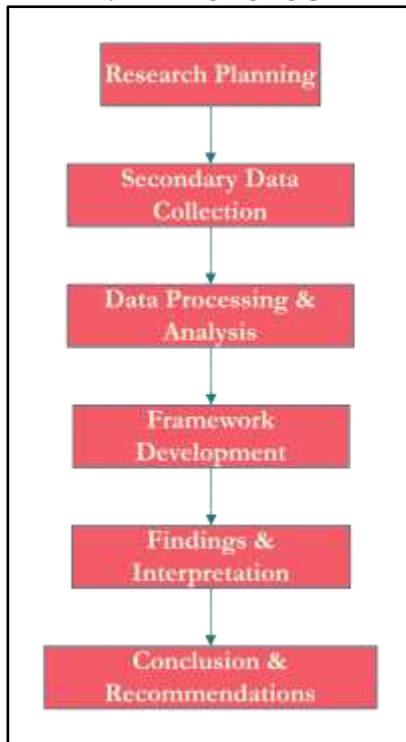


Fig 5. Research Flow Diagram

The first step of research is the research planning of the problem, aim, and objectives setting to which relevant research questions belong. Secondary data collection involves reaching out to credible sources such as journals, reports, databases and academic publications to gather relevant data is the next step. Data processing and analysis are carried out after data collection using literature screening, coding and theme analysis to identify themes and patterns within the data. The results are then fed into developing the framework, where a conceptual model of human-in-the-loop agentic workflows is created [17]. This is followed by an interpretation of the results by comparing the results with existing literature. Lastly, conclusions and suggestions for further investigations and future development of agentic AI systems are made to suggest possible avenues in which future research ought to occur in order to help inform enterprise decisions.

A. Data Collection

Data collection is the first step in the research methodology that entails obtaining information from a credible academic journal, research paper, industry report and from credible online databases when applying the relevant data. Data availability is directed by its use in agentic communication activities, Human in the Loop (HiL) environments and enterprise decision making [18]. To ensure the accuracy and validity of findings, only recent, quality and reliable

sources are used. This stage sets the groundwork for the study by ensuring that it is based on evidence and in line with existing bodies of theory and practice in AI for enterprise applications. Selecting data carefully will make sure that data remains consistent throughout the different phases of the research process, which includes doing meaningful analysis.

B. Data Preprocessing

Data preprocessing involves refining and organizing the collected secondary data to improve its quality and usability for analysis. This stage includes removing irrelevant, duplicate, or outdated information and categorizing the remaining data into key themes such as automation, human oversight, governance, risk management, and enterprise decision-making [19]. The information is structured systematically to ensure consistency across sources and to enable easier comparison. Preprocessing also involves summarizing complex literature into simplified formats without losing essential meaning. This step is critical because it enhances data reliability and ensures that only relevant and high-quality information is used for further analysis and interpretation in the research.

C. Data Analysis

The data analysis process includes a systematic analysis of secondary data in order to gain insight into the patterns, relationships and lessons learned about agentic workflows and human-in-the-loop decision-making systems. The thematic analysis approach is employed with basic statistical and comparative characteristics to support the structuring of the findings [20]. The data analysis first starts with data summarization, that is, when large sets of data or literature results are reduced to measurable indicators. The most commonly used equation is the mean or average, which is used to determine the average value of a quantitative data set:

$$\bar{x} = \frac{\sum x_i}{n}$$

where $\bar{x}$ represents the mean value, x_i represents each data point, and n is the total number of observations. This can aid in comprehending regular performance metrics like effectiveness, risk degree, or choice precision within enterprise systems.

The standard deviation is used to measure the variability and consistency of data:

$$\sigma = \frac{\sum (x_i - \bar{x})^2}{n}$$

This equation can be used to determine the difference between decisions made by AI systems and human interventions and the desired result, which serves as a measure of the system's stability and reliability.

Comparisons with results of human-validated decisions are presented as the percentage difference between the automated and human decisions.

$$\%Difference = \frac{|A - B|}{(A + B)/2} \times 100$$

where A is AI-generated outputs and B is human-validated outputs. This aids in the assessment of how well AI agents are aligned with human oversight in decisions.

Additionally, correlations between variables, like between automation level and decrease of errors, can be studied with correlation analysis:

$$r = \frac{\sum (x - \bar{x})(y - \bar{y})}{\sqrt{\sum (x - \bar{x})^2 \sum (y - \bar{y})^2}}$$

This enables one to assess the extent of the operational performance and risk mitigation impact of automation. Overall, these tools and equations help to draw a structured and quantitative interpretation on secondary data and draw reliable conclusions on the effectiveness of human-in-the-loop agents in enterprise environments.

D. System Architecture

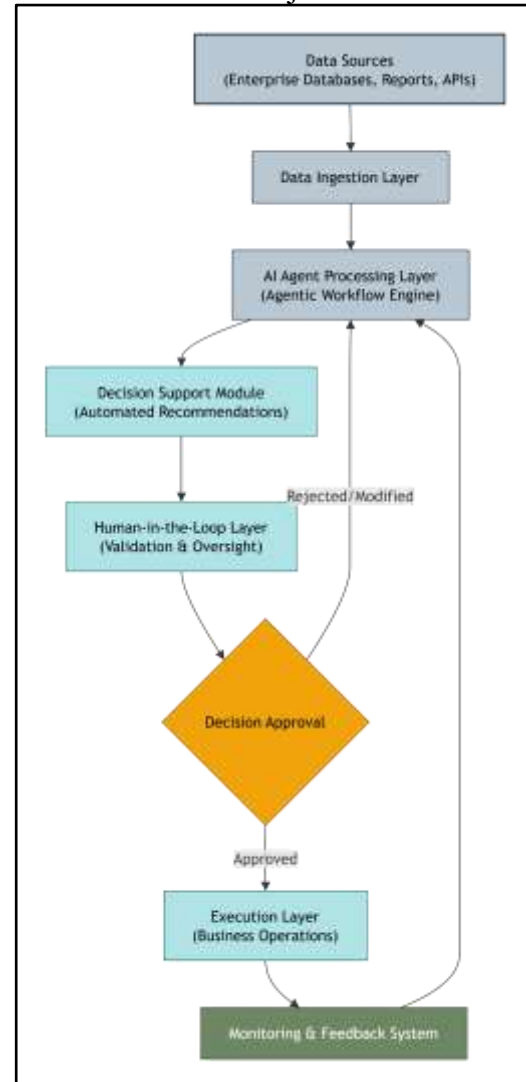


Fig 6. System Architecture text diagram

System Architecture is the structured framework in the enterprise that integrates human decision-making with AI-based agentic workflow. It lays out how information "flows" from elements of the system, from data input to AI systems, decision-making systems and human ones supervising operations. Processes are created, automated and managed according to architecture to achieve optimal process targets and compliance needs. Ethical decision-making with strategic trickle-down and human input points towards validating AI output, as well as to manage exceptions strategically [21]. Enterprise systems are scalable, flexible and reliable due to their structured nature. Overall, system architecture is a structured way of implementing successful human-in-the-loop systems, supporting in the process a balance between automation and human monitoring and decision making.

E. Ethical considerations

Ethical issues are crucial when building and implementing enterprise decision-making systems using human-in-the-loop agentic workflows. The utilization of AI agents raises some concerns regarding transparency, prejudice, responsibility, and the privacy of data [22]. Explainability of the need to make automated decisions is essential to guarantee trust and regulatory compliance. Human supervision will reduce the bias of algorithms used in the decision-making process in order to establish fair decisions [23]. Companies can protect their information that is sensitive to their security through the application of safe encryption, safe storage and safe access control. The ethical implementation of AI should also comply with legal principles such as the GDPR and good governance policies [24]. Through responsible innovation, ethical standards help create the required forms of trust in the stakeholders and sustainable enterprise processes with the help of AI without values that can be regarded as compromised.

VI. DATA ANALYSIS

Explanation of Proposed System Architecture

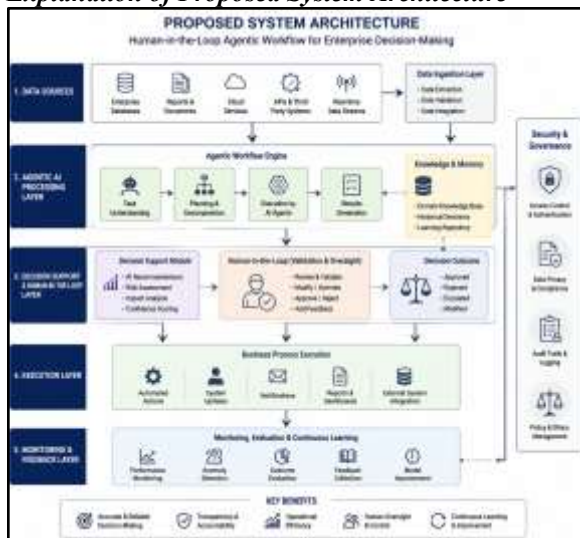


Fig 7. Proposed System Architecture for Human-in-the-Loop Agentic Workflow in Enterprise Decision-Making

The architecture of the proposed system shows a structured human-in-the-loop agentic workflow for an enterprise decision-making framework. It starts with integration from various data sources, like enterprise databases, cloud systems, and APIs to the data ingestion layer. The agentic AI processing layer involves understanding what tasks to perform, planning the steps to take, doing the tasks, and generating the results using intelligent agents [25]. The Outputs then go to the Decision Support module and are human-validated and overseen for

recommendations. The outcome of decisions that are approved is passed down to the execution stage to perform business operations [26]. The whole system is continually monitored and feedback is provided for further improvement. The security, governance, and compliance elements provide transparency, accountability, and the ethical use of AI throughout the workflow.

Data Analysis with Performance

AI-only Accuracy: 79.70%, Error Rate: 20.30%

HITL Accuracy: 93.68%, Error Rate: 6.32%

Decision Alignment with AI-only: 86.02%

Fig 8. AI-only vs HITL Accuracy Summary

The accuracy of AI-only system is 0.7970 with 20.30% error rate. Human-in-the-loop (HITL) increased the accuracy rate up to 93.68% and the number of errors down to 6.32%. Human correction significantly improves overall prediction accuracy as seen by the fact that the accuracy of the decisions made with the AI's predictions alone is 86.02% accurate.

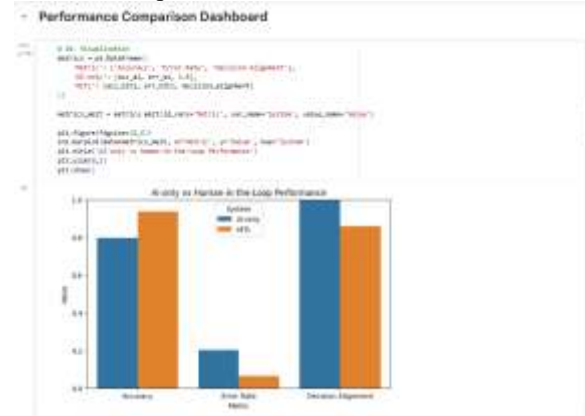
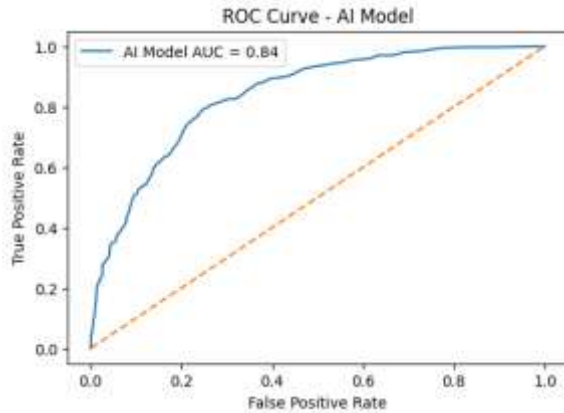


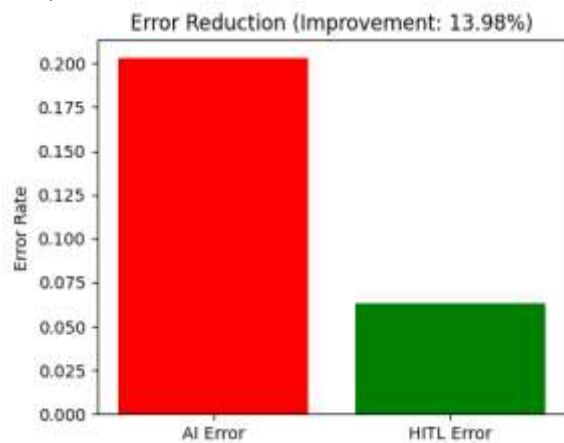
Fig 9. Performance Comparison Dashboard

The dashboard provides a visual comparison for accuracy, error rate, and decision alignment between AI models and HITL. Evaluation shows that HITL imparts higher accuracy and minimizes errors, as indicated by the alignment, meaning that the model maintains consistency with AI recommendations, further contributing to governance. The evaluation proves that HITL offers greater accuracy and minimizes errors, whereas alignment means that HITL maintains consistency with AI suggestions and contributes to governance further.

10.48047/jocaaa.2025.34.10.33

**Fig 10. ROC Curve - AI Model**

The ROC curve visualizes the performance of AI models, the True Positive Rate (TPR) against the False Positive Rate (FPR). A before-human-intervention post-hoc assessment of the baseline quality of an AI model is given by the Area Under Curve (AUC = 0.84).

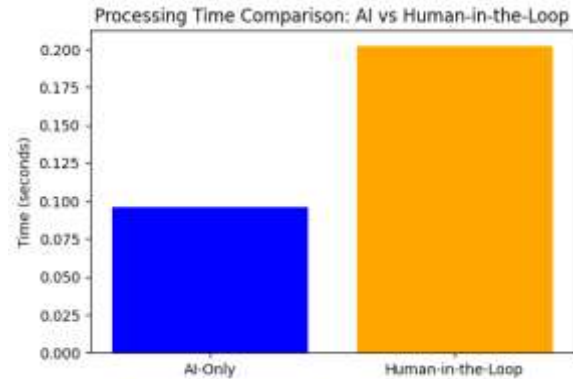
**Fig 11. Error Reduction (Improvement)**

The AI-only errors accounted for around 20% of the errors, compared to HITL errors around 6.3% which shows an increase of 13.98% in errors that are AI-only. By manually addressing AI misclassifications, human intervention diminishes operational risks and improves system reliability, which is valuable in enterprise decision-making.

	System	Accuracy	Error Rate	Decision Alignment
0	AI-Only	0.797019	0.202981	1.000000
1	Human-in-the-Loop	0.936835	0.063165	0.860185

Fig 12. Performance Summary Table

View of system metrics tabular: accuracy down from 20.3% to 6.3%, while using only AI increases it to 79.7% and using HITL increases to 93.7%. Decision alignment 86.0% indicates substantial alignment with AI with human validation, further confirming the effectiveness of HITL system.

**Fig 13. Processing Time Comparison: AI vs HITL**

The above-mentioned advantage is the speed at which the AI-only system processes the information (~ 0.10-0.12 sec), thanks to full automation. Human validation adds some delay (~0.20 sec) to HITL system, indicating trade-off between speed and decision quality resulting from human validation in enterprise operational workflows.

Discussion

The results highlight that in comparison to AI alone systems, Enterprise decision-making performance has improved from the addition of Human-in-the-Loop (HITL) mechanisms. It was found that the model without the involvement of AI has a higher error rate (20.30), in comparison with HITL model and lower error rate (6.32), which indicates greater reliability and control of HITL system. Recoded to give an even distribution of control over decisions, the percentage of aligning the decision at 86.02 instead of 70 showing that the human check enhanced AI results without losing the consistency of the model. The value 0.84 of ROC-AUC means that the AI model would be good in predicting results even before any human intervention. Moreover, the processing time analysis shows that there is a trade-off between speed and accuracy: the processing of AI-controlled equates to a short processing time but could be inaccurate. The processing of HITL is a slight delay due to the validation process but with a higher accuracy level and reliability [27]. Humans play their responsibility to improve the reliability and safety of automated decision-making systems by integrating the human-focused AI systems. AI is a duel that strengthens the argument that a hybrid AI-human systems should be more suitable when it comes to enterprise governance and usage in the context of risk [28]. The findings are evaluate human-in-the-loop agentic AI systems, which also found that the introduction of human oversight to AI systems increased the accuracy of decisions, the number of mistakes is reduced and the better the governance. This highlights greater dependability, morality and efficiency in utilizing hybrid AI-human

operating workflows in the enterprise decision-making systems.

V. CONCLUSION AND FUTURE RESEARCH

Future Aspects

In the future, there can be more advanced systems of explainable AI. AI adapts autonomously while in operation, and AI that codes with autonomous governing structuring in the research arena [29]. Being transparent, scalable, and decision intelligent will help them to achieve more accurate, ethical and efficient operations.

Conclusion

The study results indicated that the human-in-the-loop agentic workflow is important to improve enterprise decision processes, while balancing automation and human effort. The downside to the AAI-only models is that they are fast, but the hybrid models are more accurate, present fewer risks and provide better governance especially important these days for businesses.

VI. REFERENCE

- [1] Elmokadem, T. and Savkin, A.V., 2021. Towards fully autonomous UAVs: A survey. *Sensors*, 21(18), p.6223.
- [2] Frenette, J., 2023. Ensuring human oversight in high-performance AI systems: A framework for control and accountability. *World Journal of Advanced Research and Reviews*, 20(2), pp.1507-1516.
- [3] Samdani, G., Paul, K. and Saldanha, F., 2023. Agentic AI in the age of hyper-automation. *World Journal of Advanced Engineering Technology and Sciences*, 8(1), pp.416-427.
- [4] Retzlaff, C.O., Das, S., Wayllace, C., Mousavi, P., Afshari, M., Yang, T., Saranti, A., Angerschmid, A., Taylor, M.E. and Holzinger, A., 2024. Human-in-the-loop reinforcement learning: A survey and position on requirements, challenges, and opportunities. *Journal of Artificial Intelligence Research*, 79, pp.359-415.
- [5] Shneiderman, B., 2020. Bridging the gap between ethics and practice: guidelines for reliable, safe, and trustworthy human-centered AI systems. *ACM Transactions on Interactive Intelligent Systems (TiiS)*, 10(4), pp.1-31.
- [6] Herrmann, T. and Pfeiffer, S., 2023. Keeping the organization in the loop: a socio-technical extension of human-centered artificial intelligence. *Ai & Society*, 38(4), pp.1523-1542.
- [7] Grønsund, T. and Aanestad, M., 2020. Augmenting the algorithm: Emerging human-in-the-loop work configurations. *The Journal of Strategic Information Systems*, 29(2), p.101614.
- [8] Sterz, S., Baum, K., Biewer, S., Hermanns, H., Lauber-Rönsberg, A., Meinel, P. and Langer, M., 2024, June. On the quest for effectiveness in human oversight: Interdisciplinary perspectives. In *Proceedings of the 2024 ACM Conference on Fairness, Accountability, and Transparency* (pp. 2495-2507).
- [9] Kim, Y. and Xu, Y., 2024. Operational risk management: Optimal inspection policy. *Management Science*, 70(6), pp.4087-4104.
- [10] Cherukuri, R. and Yarram, V.K., 2024. From Intelligent Automation to Agentic AI: Engineering the Next Generation of Enterprise Systems. *International Journal of Emerging Research in Engineering and Technology*, 5(4), pp.142-152.
- [11] O'Grady, C.G. and O'Grady, C., 2024. Agentic workflows in the practice of law—ai agents as ethics counsel. *Arizona Legal Studies Discussion Paper*, pp.25-03.
- [12] Katipelly, A., 2024. Hierarchical Agentic Orchestration for Microservices: A Neuro-Symbolic Framework for Dynamic Workflow Composition in Decentralized Financial Systems. *International Journal of Emerging Research in Engineering and Technology*, 5(4), pp.165-174.
- [13] Parker, S.K. and Grote, G., 2022. Automation, algorithms, and beyond: Why work design matters more than ever in a digital world. *Applied psychology*, 71(4), pp.1171-1204.
- [14] Alao, O.B., Dudu, O.F., Alonge, E.O. and Eze, C.E., 2024. Automation in financial reporting: A conceptual framework for efficiency and accuracy in US corporations. *Global Journal of Advanced Research and Reviews*, 2(02), pp.040-050.
- [15] Binns, R., 2022. Human Judgment in algorithmic loops: Individual justice and automated decision-making. *Regulation & governance*, 16(1), pp.197-211.

10.48047/jocaaa.2025.34.10.33

- [16] Enemosah, A., 2021. Intelligent decision support systems for oil and gas control rooms using real-time AI inference. *International Journal of Engineering Technology Research & Management*, 5(12), pp.236-244.
- [17] Herrmann, T. and Pfeiffer, S., 2023. Keeping the organization in the loop: a socio-technical extension of human-centered artificial intelligence. *Ai & Society*, 38(4), pp.1523-1542.
- [18] Trakadas, P., Simoens, P., Gkonis, P., Sarakis, L., Angelopoulos, A., Ramallo-González, A.P., Skarmeta, A., Trochoutsos, C., Calvo, D., Pariente, T. and Chintamani, K., 2020. An artificial intelligence-based collaboration approach in industrial iot manufacturing: Key concepts, architectural extensions and potential applications. *Sensors*, 20(19), p.5480.
- [19] Odedina, E.A., 2023. Redefining governance, risk, and compliance (GRC) in the digital age: integrating AI-Driven risk management frameworks. *World Journal of Advanced Engineering Technology and Sciences*, 10(1), pp.264-282.
- [20] Talan, T., Doğan, Y. and Batdı, V., 2020. Efficiency of digital and non-digital educational games: A comparative meta-analysis and a meta-thematic analysis. *Journal of Research on Technology in Education*, 52(4), pp.474-514.
- [21] Tian, H. and Suo, D., 2021. The trickle-down effect of responsible leadership on employees' pro-environmental behaviors: Evidence from the hotel industry in China. *International Journal of Environmental Research and Public Health*, 18(21), p.11677.
- [22] Camilleri, M.A., 2024. Artificial intelligence governance: Ethical considerations and implications for social responsibility. *Expert systems*, 41(7), p.e13406.
- [23] Szafran, D. and Bach, R.L., 2024. "The human must remain the central focus": Subjective fairness perceptions in automated decision-making. *Minds and Machines*, 34(3), p.24.
- [24] Larsson, S., 2020. On the governance of artificial intelligence through ethics guidelines. *Asian Journal of Law and Society*, 7(3), pp.437-451.
- [25] Dumas, M., Fournier, F., Limonad, L., Marrella, A., Montali, M., Rehse, J.R., Accorsi, R., Calvanese, D., De Giacomo, G., Fahland, D. and Gal, A., 2023. AI-augmented business process management systems: a research manifesto. *ACM Transactions on Management Information Systems*, 14(1), pp.1-19.
- [26] Gao, S., Fang, A., Huang, Y., Giunchiglia, V., Noori, A., Schwarz, J.R., Ektefaie, Y., Kondic, J. and Zitnik, M., 2024. Empowering biomedical discovery with AI agents. *Cell*, 187(22), pp.6125-6151.
- [27] Shaikh, S.A., Memon, M. and Kim, K.S., 2021. A multi-criteria decision-making approach for ideal business location identification. *Applied sciences*, 11(11), p.4983.
- [28] Ostheimer, J., Chowdhury, S. and Iqbal, S., 2021. An alliance of humans and machines for machine learning: Hybrid intelligent systems and their design principles. *Technology in Society*, 66, p.101647.
- [29] Ng, G.W. and Leung, W.C., 2020. Strong artificial intelligence and consciousness. *Journal of Artificial Intelligence and Consciousness*, 7(01), pp.63-72.