

Toward Integrated AI-Powered Autonomous Agents for End-to-End Data Analysis and Task Automation

Hafiz Noman Javad

Southern Arkansas University
Magnolia, Arkansas, USA

hjaved5480@muleriders.saumag.edu

Praneet Kc

Southern Arkansas University
Magnolia, Arkansas, USA

pkc4990@muleriders.saumag.edu

Hayder Zhgair

Southern Arkansas University
Magnolia, Arkansas, USA

hrzghair@saumag.edu

Abstract

The rapid advancement of Large Language Models (LLMs) and autonomous Artificial Intelligence (AI) agents has transformed enterprise analytics by enabling intelligent automation, natural language interaction, and data-driven decision support. Despite these advances, existing research remains fragmented, with most studies focusing on isolated capabilities such as LLM reasoning, multi-agent collaboration, tool integration, intelligent automation, explainable AI, or cybersecurity. This fragmentation limits the development of comprehensive enterprise AI systems capable of supporting end-to-end analytical workflows. To address this gap, this study proposes a Secure and Explainable Multi-Agent Large Language Model Framework for Autonomous Enterprise Data Analytics and Decision Support. The framework is developed using the Design Science Research (DSR) methodology and integrates specialized AI agents responsible for data understanding, workflow orchestration, analytical reasoning, external tool invocation, visualization, validation, recommendation generation, security monitoring, governance, and human-in-the-loop decision support. A representative manufacturing case study demonstrates the operational workflow of the proposed architecture, illustrating how enterprise data are transformed into secure, explainable, and actionable recommendations through coordinated multi-agent collaboration. The framework is further evaluated against literature-derived design requirements, demonstrating its architectural completeness, interoperability, scalability, and alignment with enterprise AI principles. The proposed framework contributes to the literature by integrating previously disconnected research areas into a unified enterprise architecture while providing practical design guidance for developing secure, explainable, and trustworthy AI-enabled decision-support systems. Although demonstrated using a manufacturing scenario, the framework is domain-independent and can be adapted to healthcare, finance, logistics, energy systems, and other data-intensive enterprise environments.

Keywords

Large Language Models (LLMs); Multi-Agent Systems; Enterprise Artificial Intelligence; Autonomous AI Agents; Enterprise Data Analytics

1. Introduction

Artificial Intelligence (AI) has become a fundamental enabler of enterprise digital transformation by supporting organizations in extracting actionable insights from rapidly growing volumes of structured and unstructured data. Recent advances in Large Language Models (LLMs) have significantly expanded the capabilities of AI beyond traditional predictive analytics, enabling natural language interaction, code generation, reasoning, workflow automation, and decision support. These capabilities have accelerated the adoption of AI across manufacturing, healthcare, finance, logistics, cybersecurity, and other data-intensive domains, where organizations increasingly rely on intelligent systems to improve operational efficiency and support complex decision-making. The emergence of autonomous AI agents represents the next stage in this evolution. Unlike conventional LLM-based applications that primarily generate text or analytical outputs, AI agents can autonomously plan tasks, coordinate multiple analytical activities, invoke external software tools and APIs, retrieve information from enterprise databases, and collaborate with other specialized agents to accomplish complex objectives. Consequently, enterprise AI is evolving from isolated language models toward collaborative multi-agent ecosystems capable of executing complete analytical workflows with limited human intervention. Despite these advances, existing research remains fragmented. Current studies typically investigate individual capabilities such as LLM reasoning, multi-agent collaboration, intelligent automation, external tool integration, explainable AI, cybersecurity, or AI-driven visualization independently. While these contributions have advanced specific aspects of enterprise AI, they rarely address how these technologies can be systematically integrated into a unified architecture capable of supporting end-to-end enterprise data analytics and decision support. As a result, organizations often deploy disconnected AI solutions that lack interoperability, coordinated decision-making, governance, explainability, and continuous organizational learning.

This fragmentation presents a significant research gap. Enterprise decision-making is inherently multidisciplinary and requires coordinated interaction among heterogeneous data sources, analytical models, enterprise information systems, external computational tools, organizational policies, security mechanisms, and human expertise. A comprehensive enterprise AI framework should therefore provide not only intelligent reasoning but also secure orchestration, explainability, governance, validation, and human-centered collaboration throughout the complete analytical lifecycle. To address this gap, this study proposes a Secure and Explainable Multi-Agent Large Language Model Framework for Autonomous Enterprise Data Analytics and Decision Support. The framework is developed using the Design Science Research (DSR) methodology and integrates specialized AI agents responsible for data understanding, workflow orchestration, analytical reasoning, external tool invocation, visualization, validation, recommendation generation, security monitoring, and human-in-the-loop decision support. Unlike existing approaches that focus on isolated AI capabilities, the proposed framework provides a unified architectural foundation that enables coordinated interaction among multiple intelligent agents while maintaining interoperability, transparency, scalability, and organizational governance. The primary contributions of this study are fourfold. First, this research synthesizes recent literature to identify the fundamental design requirements for next-generation enterprise AI systems. Second, it proposes a secure and explainable multi-agent enterprise architecture that integrates specialized AI agents, external analytical tools, security mechanisms, explainability, governance, and human oversight into a unified decision-support framework. Third, the applicability of the proposed framework is demonstrated through a representative manufacturing case study illustrating an end-to-end enterprise analytics workflow. Finally, the proposed framework is evaluated against literature-derived design requirements, demonstrating its architectural completeness, interoperability, scalability, and applicability across multiple enterprise domains.

The remainder of this paper is organized as follows. Section 2 reviews the relevant literature on Large Language Models, multi-agent systems, intelligent automation, explainable AI, and enterprise decision support. Section 3 presents the Design Science Research methodology. Section 4 describes the proposed framework in detail. Section 5 discusses and evaluates the framework against established design requirements, and Section 6 concludes the paper with key findings, limitations, and directions for future research.

2. Literature Review

Recent advances in Artificial Intelligence (AI) have significantly transformed enterprise analytics by enabling intelligent automation, natural language interaction, and data-driven decision support. Among these developments, Large Language Models (LLMs) have emerged as powerful reasoning engines capable of understanding natural language, generating executable code, summarizing information, and assisting users throughout the data analysis lifecycle. Unlike traditional machine learning approaches that focus on narrowly defined prediction tasks, LLMs

provide generalized reasoning capabilities that support a wide range of analytical activities, including data preprocessing, feature engineering, visualization, reporting, and decision support. Recent studies demonstrate the growing role of LLMs in enterprise analytics. Hu et al. (2025) highlighted the importance of dataset-aware reasoning, showing that analytical performance varies across structured, semi-structured, and unstructured data. Similarly, Jansen et al. (2025) demonstrated that LLMs can automate substantial portions of the data science workflow through autonomous code generation and iterative self-correction, reducing the dependence on manual programming while improving accessibility for non-technical users. These studies indicate that LLMs are evolving from conversational assistants toward intelligent analytical collaborators capable of supporting enterprise decision-making. As enterprise problems become increasingly complex, researchers have shifted from single-agent architectures toward collaborative multi-agent systems. Rather than assigning all computational responsibilities to a single language model, multi-agent architectures distribute specialized tasks among independent agents coordinated through supervisory mechanisms. Ferrag et al. (2026) emphasized that distributed intelligence improves scalability, modularity, robustness, and adaptability, making multi-agent systems more suitable for enterprise environments that require simultaneous data processing, reasoning, optimization, and decision support.

Another significant advancement is the integration of LLMs with external computational tools and enterprise software. Liang et al. (2023) proposed TaskMatrix.AI, which enables language models to coordinate external APIs and computational resources to solve complex tasks. Similarly, Schick et al. (2023) introduced Toolformer, demonstrating that language models can autonomously determine when and how to invoke external tools during problem solving. These approaches significantly expand the practical capabilities of LLMs by combining natural language reasoning with specialized computational resources, allowing AI systems to perform tasks that extend beyond the knowledge encoded during model training. Parallel developments have occurred in intelligent automation and enterprise decision support. Jin et al. (2025) observed that modern automation platforms increasingly integrate LLMs, machine learning, cloud computing, workflow orchestration, and robotic process automation to improve organizational efficiency. At the same time, AI-enhanced dashboards are transforming traditional business intelligence platforms into interactive decision-support environments capable of identifying anomalies, generating explanations, and recommending actions (Sultana, 2025). The CPA Journal (2025) further demonstrated that Python-based automation substantially improves enterprise reporting and visualization by integrating data preprocessing, statistical analysis, and automated dashboard generation into unified analytical workflows. Despite these advances, enterprise deployment of AI continues to face significant challenges related to security, privacy, explainability, governance, and organizational trust. AI agents capable of interacting with external tools introduce additional risks, including prompt injection, unauthorized tool invocation, data leakage, hallucinations, and adversarial attacks. Recent studies emphasize that these challenges cannot be addressed through isolated security mechanisms but require integrated architectural approaches that incorporate validation, monitoring, governance, and explainability throughout the AI lifecycle. Human-centered AI has therefore emerged as an essential design principle, ensuring that human experts remain responsible for reviewing, validating, and approving AI-generated recommendations before organizational implementation.

Although existing research has made substantial contributions to individual aspects of enterprise AI, the literature remains fragmented. Current studies primarily focus on specific capabilities such as LLM reasoning, multi-agent collaboration, intelligent automation, external tool integration, explainable AI, cybersecurity, or visualization independently. Few studies propose a unified architectural framework that integrates these complementary technologies into a coordinated enterprise decision-support ecosystem. Consequently, organizations frequently rely on disconnected AI solutions that lack interoperability, coordinated orchestration, governance, and continuous organizational learning. Based on this literature synthesis, this study identifies eight fundamental design requirements for next-generation enterprise AI systems: dataset-aware reasoning, distributed intelligence, intelligent tool orchestration, automated analytics, integrated security and privacy, explainability, human-centered collaboration, and continuous organizational learning. These requirements provide the theoretical foundation for developing the proposed Secure and Explainable Multi-Agent LLM Framework presented in the following sections.

3. Research Methodology

This study adopts the Design Science Research (DSR) methodology because its primary objective is to develop a novel enterprise decision-support framework rather than evaluate an existing phenomenon. Design Science Research is well suited for studies that focus on designing innovative artifacts, such as models, frameworks, architectures, and methods, to address practical organizational problems while contributing to academic knowledge (Hevner et al., 2004;

Gregor & Hevner, 2013). Unlike explanatory research, which aims to understand existing systems, DSR emphasizes the systematic design, development, and evaluation of artifacts that solve identified research gaps. The proposed artifact in this study is a Secure and Explainable Multi-Agent Large Language Model Framework for Autonomous Enterprise Data Analytics and Decision Support. The framework integrates recent advances in Large Language Models (LLMs), multi-agent systems, intelligent automation, enterprise analytics, explainable AI, cybersecurity, and human-centered decision support into a unified enterprise architecture. Given the multidisciplinary nature of the problem, Design Science Research provides an appropriate methodological foundation for synthesizing knowledge from multiple research domains while producing a practical architectural solution for enterprise decision-making. The research process consists of five sequential phases. First, a comprehensive literature review was conducted to identify current developments, challenges, and research gaps related to enterprise AI, autonomous agents, intelligent automation, and decision-support systems. Second, the findings from the literature were synthesized to derive the fundamental design requirements necessary for next-generation enterprise AI systems. Third, these requirements guided the design and development of the proposed multi-agent framework, including its architectural layers, specialized AI agents, security mechanisms, governance components, and human-in-the-loop decision-support process. Fourth, the applicability of the proposed framework was demonstrated through a representative manufacturing case study that illustrates the end-to-end analytical workflow. Finally, the framework was evaluated by assessing its ability to satisfy the literature-derived design requirements and by comparing its architectural capabilities with representative approaches reported in previous studies.

This methodology ensures that every architectural component is directly supported by evidence from existing literature, providing clear traceability between the identified research gap, the derived design requirements, and the proposed framework. Such traceability enhances the rigor, transparency, and reproducibility of the research while ensuring that the proposed architecture addresses both theoretical and practical challenges associated with enterprise AI systems. Figure 1 illustrates the overall Design Science Research process adopted in this study, beginning with literature analysis, followed by requirement identification, framework development, demonstration through a representative case study, and evaluation of the proposed artifact. This structured methodology provides a systematic foundation for developing secure, explainable, and human-centered enterprise AI architectures that support end-to-end autonomous data analytics and decision-making.

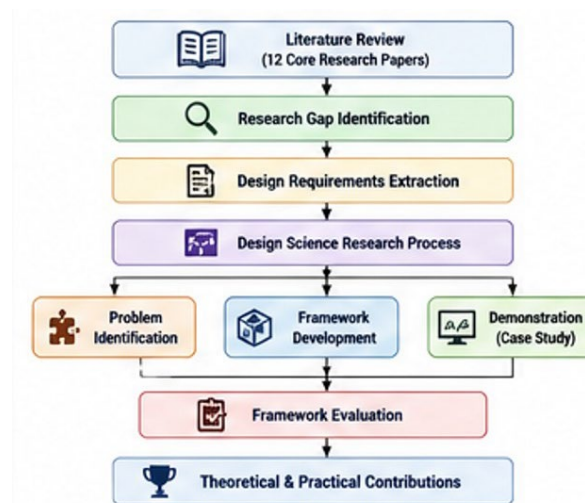


Figure 1: Research Methodology Based on Design Science Research.

The proposed methodology ensures traceability throughout the research process by linking every architectural decision to evidence derived from literature. This traceability enhances both the rigor and transparency of the study while providing a structured foundation for future implementation and empirical evaluation.

4. Framework Development

The proposed framework is the primary artifact of this research and was developed to address the fragmentation identified in existing enterprise AI literature. Rather than relying on a single Large Language Model (LLM) or isolated analytical components, the framework adopts a secure and explainable multi-agent architecture that integrates specialized AI agents, enterprise information systems, external computational tools, governance mechanisms, and human expertise into a unified decision-support ecosystem. The design follows systems engineering principles, where complex organizational problems are addressed through coordinated interaction among specialized components instead of centralized intelligence. The framework is designed around the concept of an end-to-end enterprise analytics lifecycle, transforming raw organizational data into actionable recommendations through a sequence of coordinated analytical processes. Enterprise data originating from databases, cloud platforms, IoT devices, ERP systems, CRM systems, manufacturing execution systems, and external APIs are first analyzed to determine their characteristics and quality. Based on this understanding, a central Coordinator Agent decomposes the user request into specialized analytical tasks and dynamically assigns them to the most appropriate AI agents. Each agent performs a specific responsibility, including data understanding, analytical reasoning, statistical analysis, visualization, external tool invocation, validation, recommendation generation, and explanation generation. This distributed approach improves scalability, flexibility, maintainability, and fault tolerance while enabling the framework to support increasingly complex enterprise workflows. A distinguishing feature of the proposed framework is the integration of security, explainability, and governance as core architectural components rather than optional post-processing functions. Every interaction involving data access, agent communication, external tool execution, and recommendation generation is monitored through dedicated validation and security mechanisms that ensure transparency, accountability, and organizational compliance. Human experts remain actively involved throughout the analytical process by reviewing AI-generated recommendations, providing contextual judgment, validating critical decisions, and contributing feedback that continuously improves future analytical performance. This human-in-the-loop approach balances the efficiency of autonomous AI with the experience and accountability of domain experts.

The overall architecture consists of seven interconnected layers: Enterprise Data Layer, Data Understanding Layer, AI Coordination Layer, Specialized Multi-Agent Layer, Security, Validation and Governance Layer, Decision Support and Human Interaction Layer, and Continuous Learning and Feedback Layer. Each layer performs a distinct function while collaborating with the others to ensure secure, explainable, and efficient enterprise decision support. Together, these layers establish a comprehensive AI ecosystem capable of integrating heterogeneous enterprise data, orchestrating multiple intelligent agents, coordinating external analytical tools, generating explainable recommendations, and continuously learning from organizational feedback. Unlike many existing AI-based enterprise solutions that emphasize individual technologies, the proposed framework provides a unified architectural foundation capable of integrating recent advances in Large Language Models, autonomous AI agents, intelligent automation, explainable AI, cybersecurity, and enterprise decision support. Although demonstrated using a manufacturing case study, the architecture is intentionally domain-independent and can be adapted to healthcare, finance, logistics, energy systems, smart cities, cybersecurity, and other data-intensive enterprise environments. Figure 2 illustrates the overall architecture of the proposed Secure and Explainable Multi-Agent LLM Framework, while the following subsections describe each architectural layer and specialized agent in detail.

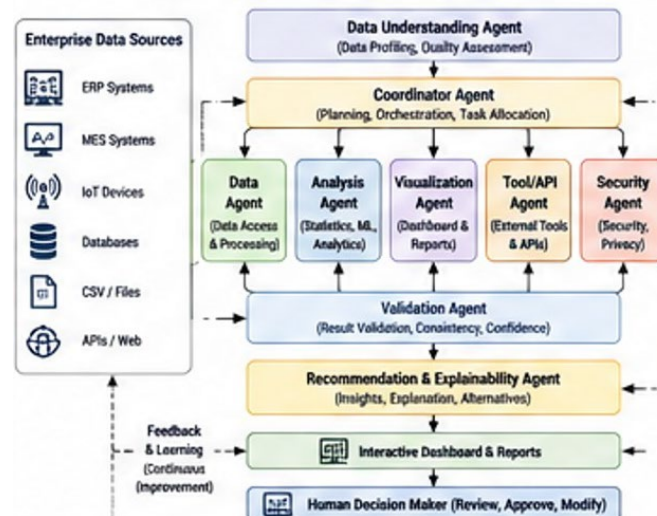


Figure 2: Proposed Secure Explainable Multi-Agent Framework.

5. Discussion

The proposed Secure and Explainable Multi-Agent LLM Framework addresses a critical gap identified in the existing literature by integrating multiple AI capabilities into a unified enterprise decision-support architecture. While previous studies have demonstrated the potential of Large Language Models, multi-agent systems, intelligent automation, tool orchestration, explainable AI, and enterprise analytics, these technologies have largely been investigated independently. The proposed framework bridges these research streams by providing a comprehensive architecture that enables coordinated collaboration among specialized AI agents, enterprise information systems, external computational tools, and human decision-makers. A key strength of the framework is its distributed intelligence. Rather than relying on a single LLM to perform all analytical tasks, computational responsibilities are assigned to specialized agents, each designed for a specific function such as data understanding, analytical reasoning, visualization, validation, security monitoring, or recommendation generation. This modular design improves scalability, flexibility, maintainability, and fault tolerance while allowing organizations to expand or modify individual components without affecting the overall architecture. Another significant contribution is the integration of security, explainability, and governance throughout the analytical workflow. Unlike many existing AI systems that address these aspects as post-processing activities, the proposed framework embeds validation, monitoring, and governance mechanisms within every stage of the decision-making process. This design enhances transparency, improves organizational trust, supports regulatory compliance, and reduces risks associated with unauthorized access, prompt injection, hallucinations, and unreliable AI-generated recommendations. The framework also emphasizes human-centered decision-making by maintaining a human-in-the-loop approach. Rather than replacing human expertise, AI-generated recommendations are presented to decision-makers for review, validation, and contextual interpretation before implementation. This collaborative approach combines the efficiency of autonomous AI with human judgment and domain expertise, resulting in more reliable and accountable enterprise decision support.

To assess the effectiveness of the proposed framework, it is evaluated against the design requirements identified through the literature review. These requirements represent the essential characteristics expected from next-generation enterprise AI systems and provide a structured basis for determining whether the proposed architecture adequately addresses the limitations of existing approaches. The evaluation presented in Table 1 demonstrates how each design requirement is incorporated into the proposed framework and highlights its overall architectural completeness and enterprise applicability.

Table 1: Evaluation of the Proposed Framework Against Identified Design Requirements.

Design Requirement	Literature Support	Framework Component	Evaluation
Dataset-aware intelligence	Hu et al. (2025)	Data Understanding Agent	Fully addressed
Distributed intelligence	Ferrag et al. (2026)	Multi-Agent Architecture	Fully addressed
Tool orchestration	Liang et al. (2023); Schick et al. (2023)	Tool/API Agent	Fully addressed
Automated analytics	Jansen et al. (2025)	Analysis Agent	Fully addressed
Intelligent visualization	Sultana (2025); <i>The CPA Journal</i> (2025)	Visualization Agent	Fully addressed
Security & privacy	Berini et al. (2026); Tang et al. (2025)	Security Layer	Fully addressed
Human-centered decision support	Guyton et al. (2025)	Human Decision Layer	Fully addressed
Continuous learning	Literature synthesis	Feedback Layer	Fully addressed

The evaluation demonstrates that the proposed Secure and Explainable Multi-Agent LLM Framework successfully satisfies all eight design requirements identified through the literature review. By explicitly mapping each requirement to a corresponding architectural component, the framework establishes clear traceability between the research gap, the proposed solution, and the underlying theoretical foundation. This traceability strengthens the rigor of the proposed architecture and demonstrates that its design decisions are systematically derived from existing research rather than based solely on conceptual assumptions. Compared with existing enterprise AI approaches, the proposed framework adopts a systems-oriented perspective by integrating multiple complementary technologies into a single decision-support architecture. Rather than focusing on isolated capabilities such as LLM reasoning, intelligent automation, or multi-agent collaboration, the framework combines dataset-aware intelligence, distributed AI agents, external tool orchestration, automated analytics, security, explainability, governance, and human-centered collaboration within a unified enterprise ecosystem. This integration addresses the fragmentation identified in previous studies and provides a more comprehensive solution for enterprise decision-making. Another important contribution of the framework is its emphasis on modularity and interoperability. The distributed multi-agent architecture enables specialized agents to perform independent analytical tasks while coordinating through a central orchestration mechanism. Such modularity improves scalability, maintainability, and flexibility, allowing organizations to extend or replace individual components without redesigning the entire system. Furthermore, the framework supports seamless integration with existing enterprise information systems, cloud platforms, external APIs, and analytical tools, making it adaptable to diverse organizational environments.

The framework also prioritizes trustworthiness by embedding security, validation, and explainability throughout the analytical workflow. Instead of treating these features as post-processing activities, dedicated architectural components continuously monitor data access, agent communication, external tool execution, and recommendation generation. This integrated approach enhances transparency, supports regulatory compliance, and increases organizational confidence in AI-assisted decision-making. In addition, the human-in-the-loop mechanism ensures that final decisions remain under human supervision, allowing domain experts to validate AI-generated recommendations before implementation. Despite these contributions, several limitations should be acknowledged. The current study evaluates the framework conceptually through literature-derived design requirements and a representative case study rather than through prototype implementation or large-scale organizational deployment. Consequently, the computational performance, scalability, response time, user acceptance, and real-world organizational impact of the framework remain subjects for future investigation. Developing and evaluating a functional prototype in real enterprise environments would provide stronger empirical evidence of the framework's effectiveness.

Overall, the proposed Secure and Explainable Multi-Agent LLM Framework provides a comprehensive architectural foundation for next-generation enterprise AI systems. By integrating recent advances in Large Language Models, autonomous AI agents, intelligent automation, enterprise analytics, security, explainability, and human-centered decision support, the framework contributes both theoretical insights and practical guidance for developing secure, scalable, and trustworthy enterprise decision-support systems across multiple application domains.

6. Conclusion

The rapid advancement of Large Language Models (LLMs) and autonomous AI agents has created new opportunities for transforming enterprise data analytics and decision support. However, existing research has largely focused on

individual technologies rather than developing integrated enterprise architectures capable of coordinating multiple intelligent components within secure, explainable, and human-centered environments. This study addressed this research gap by proposing a Secure and Explainable Multi-Agent Large Language Model Framework for Autonomous Enterprise Data Analytics and Decision Support, developed using the Design Science Research (DSR) methodology. The proposed framework integrates specialized AI agents, enterprise information systems, external computational tools, intelligent automation, security mechanisms, explainability, governance, and human oversight into a unified architectural ecosystem. By distributing analytical responsibilities across multiple coordinated agents, the framework improves scalability, interoperability, transparency, and adaptability while supporting end-to-end enterprise decision-making. A representative manufacturing case study demonstrated the practical applicability of the proposed architecture, and the evaluation confirmed that the framework satisfies the key design requirements identified through the literature review. The primary contribution of this research is the development of a comprehensive enterprise AI architecture that bridges previously disconnected research areas, including LLM-based reasoning, multi-agent systems, tool orchestration, intelligent automation, explainable AI, cybersecurity, and human-centered decision support. Unlike existing approaches that emphasize isolated capabilities, the proposed framework provides a unified and extensible foundation that can be adapted to various enterprise domains, including manufacturing, healthcare, finance, logistics, energy systems, and smart cities.

Despite these contributions, this study has several limitations. The proposed framework has been evaluated conceptually through literature-derived design requirements and a representative case study rather than through prototype implementation and large-scale organizational deployment. Consequently, future research should focus on developing a functional prototype, conducting experimental validation in real-world enterprise environments, and evaluating computational performance, scalability, response time, user acceptance, and decision quality. Future work may also investigate advanced agent communication protocols, retrieval-augmented generation (RAG), long-term memory management, multimodal reasoning, and adaptive learning mechanisms to further enhance the capabilities of enterprise AI systems. Overall, this research provides a systematic architectural foundation for the next generation of enterprise AI. By integrating secure, explainable, collaborative, and human-centered intelligence within a unified multi-agent framework, the proposed approach contributes both theoretical insights and practical guidance for designing trustworthy AI-enabled decision-support systems capable of addressing increasingly complex organizational challenges.

References

- Arrieta, A. B., Díaz-Rodríguez, N., Del Ser, J., Benetot, A., Tabik, S., Barbado, A., García, S., Gil-López, S., Molina, D., Benjamins, R., Chatila, R., & Herrera, F. Explainable artificial intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information Fusion*, 58, 82–115, 2020, <https://doi.org/10.1016/j.inffus.2019.12.012>
- Brown, T. B., Mann, B., Ryder, N., Subbiah, M., Kaplan, J., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., Agarwal, S., Herbert-Voss, A., Krueger, G., Henighan, T., Child, R., Ramesh, A., Ziegler, D. M., Wu, J., Winter, C., ... Amodei, D. Language models are few-shot learners. *Advances in Neural Information Processing Systems*, 33, 1877–1901, 2020.
- Brynjolfsson, E., & McAfee, A. The business of artificial intelligence. *Harvard Business Review*, 2017, <https://hbr.org/2017/07/the-business-of-artificial-intelligence>
- Davenport, T. H., & Ronanki, R. Artificial intelligence for the real world. *Harvard Business Review*, 96(1), 108–116, 2018.
- Gregor, S., & Hevner, A. R. Positioning and presenting design science research for maximum impact. *MIS Quarterly*, 37(2), 337–355, 2013, <https://doi.org/10.25300/MISQ/2013/37.2.01>
- Gunning, D. Explainable artificial intelligence (XAI). Defense Advanced Research Projects Agency (DARPA), 2017, <https://www.darpa.mil/program/explainable-artificial-intelligence>
- Hevner, A. R., March, S. T., Park, J., & Ram, S. Design science in information systems research. *MIS Quarterly*, 28(1), 75–105, 2004, <https://doi.org/10.2307/25148625>
- Jennings, N. R. On agent-based software engineering. *Artificial Intelligence*, 117(2), 277–296, 2000, [https://doi.org/10.1016/S0004-3702\(99\)00107-1](https://doi.org/10.1016/S0004-3702(99)00107-1)
- Liang, Y., Ge, C., Tong, Z., Wang, Y., Xie, Y., Shan, Y., & Luo, P. TaskMatrix.AI: Completing tasks by connecting foundation models with millions of APIs (arXiv:2303.16434). *arXiv*, 2023, <https://arxiv.org/abs/2303.16434>
- OpenAI. GPT-4 technical report. *arXiv*, 2023, <https://arxiv.org/abs/2303.08774>

- Peffer, K., Tuunanen, T., Rothenberger, M. A., & Chatterjee, S. A design science research methodology for information systems research. *Journal of Management Information Systems*, 24(3), 45–77, 2007, <https://doi.org/10.2753/MIS0742-1222240302>
- Power, D. J. *Decision support systems: Concepts and resources for managers*. Quorum Books, 2002.
- Schick, T., Dwivedi-Yu, J., Dessì, R., Raileanu, R., Lomeli, M., Hambro, E., Zettlemoyer, L., Cancedda, N., & Scialom, T. Toolformer: Language models can teach themselves to use tools. *Nature*, 632, 563–571, 2023, <https://doi.org/10.1038/s41586-023-06647-8>
- Turban, E., Sharda, R., & Delen, D. *Decision support and business intelligence systems* (11th ed.). Pearson, 2018.
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., & Polosukhin, I. Attention is all you need. *Advances in Neural Information Processing Systems*, 30, 5998–6008, 2017.
- Wooldridge, M. *An introduction to multiagent systems* (2nd ed.). John Wiley & Sons, 2009.

Biographies

Hafiz Noman Javed is a graduate student in Data Science at Southern Arkansas University (SAU), with a strong academic foundation in Computer Science. His academic and professional interests lie at the intersection of Data Science, artificial intelligence, and workforce readiness, with a particular focus on applying data-driven approaches to real-world industry challenges. He has been actively involved in research aimed at bridging the gap between academic learning and industry requirements. His previous work includes contributions to projects exploring AI-driven solutions and data-centric decision-making systems. Currently, he is engaged in research that examines the alignment between STEM education and evolving workforce expectations, emphasizing how data and emerging technologies can enhance decision-making and improve industry preparedness. In addition to his research pursuits, Noman brings practical industry experience, having worked in remote roles involving data analysis, business operations, and digital systems. He is particularly interested in leveraging AI, data analytics, and automation to improve organizational efficiency and strategic outcomes. Through his academic and professional journey, he aims to contribute to the development of intelligent systems and frameworks that support better decision-making across industries.

Dr. Hayder Zghair is an Assistant Professor of Industrial Engineering and the Director of Industrial Engineering Development at Southern Arkansas University (SAU). He earned his Ph.D. and M.S. in Manufacturing Systems Engineering from Lawrence Technological University, alongside a B.S. and M.S. in Industrial and Production Engineering from the University of Technology. Before joining SAU, he served as an Assistant Teaching Professor at Pennsylvania State University and a lecturer at Kettering University. Dr. Zghair's research encompasses Flexible-Automated Manufacturing, Robotics, Sustainability Engineering, Analytical Modeling, Simulation, Optimization, and Industry 4.0/5.0 technologies like AI and IIoT. Currently, he is collaborating with Sana Bhimani and Seema Powar to investigate the STEM skills gap and to develop analytical frameworks to align academic curricula with entry-level workforce demands. An active leader in the engineering community, he serves on the Board of Directors for the ASEE Environmental Engineering Division and is an Affiliate Researcher at the Penn State Institute of Energy and the Environment. He holds professional memberships in ASEE, INFORMS, SAE, and IEOM, and actively mentors students as the faculty advisor for the SAE, SAU Baja, and IEOM student chapters at SAU.