

Journal of Integrated Science, Technology and Management

Review Article

LLM-Powered Enterprise Applications: A Systematic Review of Cloud Computing, Payment Gateway, and Risk Management Architectures

Agbilu T Oliver*

Massachusetts Institute of Technology, Chestnut Hill, Massachusetts, U.S.A

ARTICLE INFO

Article history:

Received: 11 September, 2025
Revised: 29 September, 2025
Accepted: 16 October, 2025
Published: 25 December, 2025

Keywords:

Large language models, Enterprise applications, Cloud computing, payment gateway, Fraud detection, Risk management, Event-driven architecture, Enterprise architecture

ABSTRACT

Large language models (LLMs) are increasingly embedded within enterprise technology stacks, extending beyond conversational interfaces into core operational systems including cloud infrastructure, payment processing, and risk management. This systematic review synthesizes current architectural approaches to LLM-powered enterprise applications across three interdependent layers: the cloud-computing infrastructure required to deploy LLMs reliably at production scale, payment-gateway applications in which LLMs are increasingly used for fraud detection and transaction-risk analysis (Cao et al., 2024), and the risk-management and compliance architecture required to govern LLM-driven enterprise decisions. The review draws on applied research addressing zero-downtime AI model updates in real-time inference systems (Sannidhanam, 2023a), AI-driven failure prediction and automated recovery in self-healing distributed systems (Sannidhanam, 2023b), event-driven and event-streaming architectures for high-volume transaction and claims processing (Event Streaming Architectures for High-Volume Transaction Processing, 2022; Sannidhanam, 2021), and enterprise risk-governance architecture, including configurable workflow-based risk management (Basireddy, 2022) and audit-ready compliance platforms specifically designed for LLM-regulated enterprises (Basireddy, 2023). The review finds that reliable enterprise LLM deployment depends on a layered architecture in which real-time, event-driven data infrastructure, resilient and self-healing cloud operations, and auditable governance controls are treated as co-equal design requirements alongside model performance, and it identifies payment-gateway fraud detection as one of the most mature and rapidly growing applied use cases for enterprise LLM deployment. The review concludes by identifying gaps in the integration of these three architectural layers and outlining priorities for future research on production-grade, governed LLM enterprise systems.

INTRODUCTION

Large language models (LLMs) have moved rapidly from standalone conversational applications into deeply embedded components of enterprise technology stacks, supporting functions ranging from customer service and content generation to fraud detection, transaction processing, and regulatory compliance. This expansion has been accompanied by rapid market growth: the global enterprise LLM market was valued at approximately 6.7 billion US dollars in 2024 and is projected to reach 71.1 billion US dollars by 2034, reflecting the scale at which organizations are integrating these systems into core

operations. As LLMs move from experimental pilots into production systems supporting financial transactions and risk-sensitive decisions, the architectural requirements for deploying them reliably — spanning cloud infrastructure, real-time data processing, and governance controls — have become as consequential to enterprise outcomes as the underlying model capabilities themselves.

This systematic review examines LLM-powered enterprise applications across three interdependent architectural layers named in its title: cloud computing infrastructure (Section 2), payment-gateway applications (Section 3), and risk-management architecture (Section

*Corresponding Author: Author Name

Address: Affiliation

Email ✉: email

Relevant conflicts of interest/financial disclosures: The authors declare that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.

© 2025, First author, This is an open access journal, and articles are distributed under the terms of the Creative Commons Attribution-NonCommercial-ShareAlike 4.0 License, which allows others to remix, tweak, and build upon the work non-commercially, as long as appropriate credit is given and the new creations are licensed under the identical terms.

4), followed by an integrated discussion of how these layers interact in production enterprise systems (Section 5). The review draws on applied research addressing model lifecycle management and infrastructure resilience (Sannidhanam, 2023a, 2023b), real-time event-driven data processing (Event Streaming Architectures for High-Volume Transaction Processing, 2022; Sannidhanam, 2021), LLM applications in financial services and fraud detection (Cao et al., 2024), and enterprise risk-governance architecture (Basireddy, 2022, 2023).

Review Methodology

This review synthesizes applied enterprise-architecture research, peer-reviewed academic literature on LLM applications in financial services, and industry market analysis relevant to enterprise LLM deployment. Sources were organized around the three architectural layers named in the review's title — cloud computing, payment gateways, and risk management — recognizing that production LLM-powered enterprise applications typically require all three layers to function together rather than as independent technical decisions. Where direct academic literature specific to a named architectural requirement was limited, the review draws on closely related applied research addressing the same underlying technical challenge (for example, distributed-systems resilience research applied to the specific case of LLM inference infrastructure).

Cloud Computing Infrastructure for LLM Deployment

Zero-Downtime Model Updates

A foundational operational requirement for enterprise LLM deployment is the ability to update deployed models without interrupting live inference traffic. Sannidhanam (2023a) addresses this requirement directly, describing architectural patterns for zero-downtime AI model updates in real-time inference systems, enabling organizations to deploy improved or retrained models, including updated LLM versions, without service interruption — a capability of particular importance for enterprise applications such as payment-gateway fraud detection, where inference availability directly affects transaction processing continuity.

Self-Healing Infrastructure

Beyond model updates, the underlying distributed infrastructure supporting LLM inference must be resilient to failure, given the operational consequences of inference downtime in production enterprise contexts. Sannidhanam (2023b) describes self-healing distributed systems that use AI-driven failure prediction to anticipate infrastructure failures before they occur and automated recovery mechanisms to remediate them with minimal human intervention. Applied to LLM-powered enterprise

applications, this self-healing capability addresses a critical operational gap: LLM inference infrastructure, like other distributed systems, is subject to hardware failures, network partitions, and resource-contention events that can degrade service availability unless proactively detected and remediated.

Real-Time, Event-Driven Data Processing

Many enterprise LLM applications — particularly those supporting payment processing and risk management — require processing of high-volume transactional data with minimal latency, making the underlying data-ingestion architecture a critical determinant of application performance. Applied research on event streaming architectures for high-volume transaction processing (2022) describes architectural patterns for ingesting and processing transaction data as continuous event streams, directly supporting the low-latency data availability that LLM-based real-time fraud-detection and risk-scoring applications require. Complementing this, Sannidhanam (2021) demonstrates real-time claims processing using event-driven architectures in an insurance context, illustrating how event-driven design patterns support low-latency processing of high-volume, risk-relevant transactional events in a domain closely analogous to payment-gateway transaction processing.

Autonomous Infrastructure Operations

Extending the resilience capabilities described above, more recent applied research describes autonomous AI agents deployed for cloud infrastructure operations, capable of continuously monitoring infrastructure state and executing remediation actions directly, with limited human intervention (Sannidhanam, 2025). Applied to LLM-powered enterprise applications, this represents a further maturation of the self-healing infrastructure concept introduced in Section 3.2, extending it from failure prediction and automated recovery toward fully autonomous, self-directed infrastructure operations supporting the continuous availability that production LLM inference workloads require.

Payment Gateway Applications

Payment-gateway systems represent one of the most mature and rapidly growing applied use cases for enterprise LLM deployment, driven by the natural fit between LLMs' natural-language and pattern-recognition capabilities and the fraud-detection and transaction-risk-analysis tasks central to payment processing. Cao, Li, Katsikis, Khan, He, Liu, Zhang, and Peng (2024) examine the transformative influence of LLMs across the financial sector, finding substantial application in customer support and fraud detection, where LLMs' capacity to process and contextualize large volumes of transaction and communication data enables more nuanced fraud identification than earlier rule-based or narrower

machine-learning approaches. Industry evidence corroborates this academic finding: major financial institutions have moved from pilot to enterprise-wide LLM deployment for functions including transaction analysis and fraud detection, with AI-driven fraud-detection systems at large financial institutions now processing transactions at substantial scale in real time.

Within payment-gateway applications specifically, LLMs offer capabilities beyond traditional transaction-pattern anomaly detection, including natural-language analysis of transaction metadata, merchant descriptions, and customer communications to identify fraud indicators that purely numerical models may miss. This capability is particularly relevant to detecting sophisticated fraud patterns involving social engineering or communication-based deception, where the fraudulent signal is embedded in language rather than transaction amount or frequency alone. Realizing this capability in production payment-gateway systems, however, depends directly on the cloud infrastructure requirements discussed in Section 2: fraud detection in payment processing is a latency-sensitive application, where LLM inference must complete within the transaction-authorization window, making zero-downtime model availability and event-driven data ingestion direct technical preconditions for effective payment-gateway LLM deployment, not separate architectural concerns.

Risk Management Architectures

Workflow-Embedded Risk Controls

As LLMs take on a greater role in payment and transaction-risk decisions, the ability to embed risk-management logic directly into the operational workflows where these decisions occur becomes essential. Basireddy (2022)

proposes a configurable enterprise workflow architecture that digitizes risk management by embedding risk-control logic directly into business-process workflows, allowing an LLM-flagged transaction or fraud indicator to trigger a defined operational response — additional review, transaction hold, or escalation — as part of normal workflow execution rather than through a separate, disconnected alerting system.

Audit-Ready Compliance for LLM-Regulated Enterprises

Payment and financial-services applications of LLMs operate within some of the most heavily regulated enterprise environments, making auditability of LLM-driven decisions a core architectural requirement rather than an optional enhancement. Basireddy (2023) proposes an audit-ready compliance-platform architecture specifically designed for enterprises operating LLMs in regulated environments, ensuring that LLM-driven risk decisions — such as a fraud flag or a transaction risk score generated by an LLM-based system — can be traced, reviewed, and defended to regulators and auditors after the fact. This capability directly addresses a governance gap otherwise created by LLMs' nondeterministic outputs: without a structured audit architecture, organizations may struggle to reconstruct why a specific LLM-driven risk decision was made, a significant liability in regulated payment and financial contexts.

Evidence Snapshot: Architectural Layers and Requirements

Table 1 summarizes the architectural layers, requirements, and representative sources reviewed across this article.

Integrating the Three Architectural Layers

A central finding of this review is that cloud computing, payment-gateway, and risk-management architecture

Table 1: Architectural layers and requirements for LLM-powered enterprise applications

<i>Architectural layer</i>	<i>Key requirement</i>	<i>Representative source</i>
Model lifecycle management	Zero-downtime updates to deployed AI models in live inference systems	Sannidhanam, 2023a
Infrastructure resilience	AI-driven failure prediction and automated recovery in distributed systems	Sannidhanam, 2023b
Real-time data processing	Event-driven, low-latency processing of high-volume transactions and claims	Event Streaming Architectures for High-Volume Transaction Processing, 2022; Sannidhanam, 2021
Payment-gateway application	LLM-based fraud detection and transaction-risk analysis in financial services	Cao et al., 2024
Operational risk control	Risk-control logic embedded directly into configurable enterprise workflows	Basireddy, 2022
Regulatory compliance	Auditable, traceable record of LLM-driven enterprise decisions	Basireddy, 2023
Autonomous operations	Self-directed monitoring and remediation of cloud infrastructure supporting LLM workloads	Sannidhanam, 2025

are not independent technical decisions in LLM-powered enterprise applications but interdependent layers that jointly determine whether an LLM-based application functions reliably and defensibly in production. A payment-gateway fraud-detection application built on an LLM, for example, depends on zero-downtime model-update infrastructure and event-driven data processing (Section 2) to function with acceptable latency and availability, on LLM-specific application design to deliver fraud-detection value (Section 3), and on workflow-embedded and audit-ready risk-governance architecture (Section 4) to ensure that its outputs are both operationally actionable and regulatorily defensible. Applied research reviewed here — spanning Sannidhanam's (2023a, 2023b) infrastructure-resilience work, Basireddy's (2022, 2023) governance-architecture work, and event-driven processing research (Event Streaming Architectures for High-Volume Transaction Processing, 2022; Sannidhanam, 2021) — collectively suggests that treating any one of these three layers in isolation is likely to produce an LLM-powered enterprise application that is either operationally unreliable, commercially ineffective, or regulatorily indefensible, even if the underlying LLM itself performs well in isolated testing.

Research Gaps and Future Directions

- Empirical, cross-industry evaluation of integrated cloud-payment-risk architecture for LLM applications remains limited; most reviewed sources address one architectural layer in depth rather than measuring integrated system performance across all three.
- Latency and cost trade-offs specific to LLM inference in latency-sensitive payment-authorization contexts merit further dedicated research, given the distinctive computational profile of LLM inference relative to earlier, lighter-weight fraud-detection models.
- Standardized benchmarks for evaluating LLM-based fraud-detection accuracy against traditional machine-learning baselines in production payment-gateway settings remain underdeveloped in the academic literature (Cao et al., 2024).
- The relationship between self-healing infrastructure capabilities (Sannidhanam, 2023b) and measured reductions in payment-processing downtime or fraud-detection service-level performance remains an open empirical question.
- Further research connecting workflow-embedded governance architecture (Basireddy, 2022) and audit-ready compliance systems (Basireddy, 2023) directly to regulatory examination outcomes in payment and financial-services contexts would strengthen the practical evidence base for these architectural approaches.

CONCLUSION

LLM-powered enterprise applications, particularly in payment-gateway fraud detection and transaction-risk analysis, represent one of the fastest-growing and most consequential areas of enterprise AI deployment. This systematic review demonstrates that realizing reliable, governed value from these applications depends on treating cloud-computing infrastructure, payment-gateway application design, and risk-management architecture as a single, integrated system rather than as separable technical concerns. Zero-downtime model-update capability, self-healing infrastructure, and event-driven real-time data processing provide the operational foundation on which LLM-based payment applications depend; workflow-embedded risk controls and audit-ready compliance architecture provide the governance layer that makes these applications defensible within regulated financial environments. As enterprise LLM adoption continues to accelerate, particularly within the payment and financial-services sector, future research that evaluates these architectural layers as an integrated whole, rather than in isolation, is likely to offer the most practically useful guidance for organizations seeking to deploy LLM-powered enterprise applications reliably and responsibly at scale.

REFERENCES

1. Basireddy, S. R. (2022). Digitizing risk management through configurable enterprise workflow architecture. *SAMRIDDHI: A Journal of Physical Sciences, Engineering and Technology*, 14(4), 231-239.
2. Basireddy, S. (2023). Audit-ready compliance platform architecture for large language model regulated enterprises. *SAMRIDDHI: A Journal of Physical Sciences, Engineering and Technology*, 15(01), 226-232. <https://doi.org/10.18090/samriddhi.v15i01.40>
3. Cao, X., Li, S., Katsikis, V., Khan, A. T., He, H., Liu, Z., Zhang, L., & Peng, C. (2024). Empowering financial futures: Large language models in the modern financial landscape. *EAI Endorsed Transactions on AI and Robotics*, 3(1). <https://doi.org/10.4108/airo.6117>
4. Event Streaming Architectures for High-Volume Transaction Processing. (2022). *International Journal of Advance Industrial Engineering*, 10(01), 1-8.
5. Sannidhanam, A. H. (2021). Real-time claims processing using event-driven architectures. *International Journal of Technology, Management and Humanities*, 7(01), 36-50. <https://doi.org/10.21590/07.01.02>
6. Sannidhanam, A. H. (2023a). Zero-downtime AI model updates in real-time inference systems. *International Journal of Engineering Science & Humanities*, 13(2), 70-78.
7. Sannidhanam, A. H. (2023b). Self-healing distributed systems: AI-driven failure prediction and automated recovery. *International Journal of Research & Technology*, 11(4), 168-174.
8. Sannidhanam, A. H. (2025). Autonomous AI agents for cloud infrastructure operations. *Journal of Contemporary Science and Technology Management*, 1(01), 83-100.

HOW TO CITE THIS ARTICLE: Oliver, A.T. (2025). LLM-Powered Enterprise Applications: A Systematic Review of Cloud Computing, Payment Gateway, and Risk Management Architectures. *Journal of Integrated Science, Technology and Management*, 1(1), 31-34.