

Designing Production-Grade Agentic AI for Continuous Payment Exception Monitoring, Hierarchical Authorization, and Governance in Financial Systems: An MCP, LLM, and RAG Implementation Architecture

Kedarnath Goud Kothinti

Independent Researcher, Charlotte, NC, USA

Abstract: Autonomous AI systems for financial payment operations have matured from theoretical governance proposals to deployable production architectures, yet the published literature has not kept pace with the tooling ecosystem that makes such systems practically viable. This paper advances the state of the art in two directions. First, it formalizes a production-grade agentic AI framework for continuous payment exception monitoring and hierarchical authorization, grounding the governance model in a parent-child agent architecture that coordinates risk assessment, revenue impact scoring, and customer relationship evaluation before any exception resolution action is taken. Second — and as the primary new contribution of this revision — it specifies a concrete implementation architecture in which the agentic AI system operates through Model Context Protocol (MCP) server/client connections, a large language model (LLM) reasoning core, and a Retrieval-Augmented Generation (RAG) pipeline that grounds every autonomous decision in a continuously refreshed corpus of regulatory policy, internal governance rules, and historical resolution precedents. The MCP-mediated architecture allows the agentic AI system to query payment data, compliance documents, customer relationship records, and revenue analytics through standardized tool interfaces — without requiring custom integration code for each data source. The RAG pipeline ensures that the LLM reasoning core grounds its exception analysis in specific, citable regulatory text rather than parametric knowledge, producing verifiable and audit-ready reasoning chains. A production deployment on a payment processing platform handling 2.3 million monthly transactions demonstrates a 99.7% reduction in mean time to resolve auto-eligible exceptions, an 89% ACH recovery rate, a 43% reduction in analyst workload, and zero compliance audit findings over a 12-month post-deployment window. These results demonstrate that MCP-mediated, RAG-grounded agentic AI is not a research prototype capability — it is a deployable production architecture for the financial services industry today.

Keywords: agentic AI, Model Context Protocol, large language model, retrieval-augmented generation, payment exception monitoring, hierarchical authorization, financial governance, autonomous payment systems

1. Introduction

Payment exception management — the identification, classification, and resolution of failed, disputed, or anomalous payment events — is among the highest-volume operational workflows in financial technology platforms. A mid-size FinTech processor handling two million monthly transactions will encounter between 14,000 and 35,000 exception events per month across ACH failures, gateway timeouts, duplicate submissions, and settlement reconciliation discrepancies. Managing this volume through human analyst review is neither scalable nor financially sustainable: at the analyst productivity rates documented in this paper's deployment baseline, the human-only resolution model required 68% of operational analyst hours for routine exception triage, leaving fewer than one-third of analyst capacity available for complex investigations, governance oversight, and process improvement.



The response to this operational challenge in both industry and research has been a movement toward autonomous AI systems for exception management. Ali et al. [1] frame this transition as a management innovation that reconstitutes organizational decision rights — not merely an efficiency automation. Singh et al. [2] document the governance implications of delegating financial decisions to AI agents in banking contexts. Vuković et al. [3] examine the organizational change dimensions of AI-enabled payment operations across European financial institutions. These contributions establish the theoretical case for agentic AI in payment governance. What they do not provide is a practical architecture: a concrete specification of how an agentic AI system can be built using available tooling, how it grounds its reasoning in regulatory policy without hallucination, how it integrates with diverse payment data sources without custom integration code, and how it maintains the governance accountability required for regulated financial operations.

This paper fills that gap. It introduces three architectural components that together constitute a production-deployable agentic AI system for payment exception management: (1) Model Context Protocol (MCP) server/client connections that provide the agentic AI with standardized, tool-mediated access to payment data, compliance documents, customer records, and revenue analytics; (2) an LLM reasoning core that interprets exception context and selects resolution actions within a parent-child authorization hierarchy; and (3) a RAG pipeline that grounds every reasoning step in a continuously indexed corpus of NACHA rules, PCI DSS controls, internal policy documents, and historical exception resolution precedents. The integration of these three components produces an agentic system that is simultaneously autonomous, transparent, and auditable—addressing the accountability gap that has limited the adoption of AI in regulated payment operations [14].

The paper is organized as follows. Section 2 describes the payment exception landscape and defines the four exception categories addressed by the system. Section 3 presents the overall agentic AI system architecture. Section 4 details the MCP server/client implementation and the tool registry. Section 5 describes the LLM reasoning core and parent-child authorization model. Section 6 presents the RAG pipeline and compliance document indexing methodology. Section 7 documents governance controls and regulatory alignment. Section 8 presents performance results from the production deployment. Section 9 discusses implementation considerations and limitations. Section 10 concludes with directions for future research [15].

2. Payment Exception Landscape

Payment exceptions arise at four principal points in the payment processing lifecycle: bank transfer and ACH processing, payment gateway execution, duplicate transaction detection, and settlement reconciliation. Table 1 provides a taxonomy of exception categories, exception types within each category, detection methods, and resolution pathways applicable to each.

Table 1: Payment Exception Type Classification and Resolution Pathway Taxonomy

Exception Category	Exception Type	Detection Method	Resolution Pathway
ACH / Bank Transfer	Insufficient funds, account closed, invalid routing	Return code parsing, real-time balance pre-check	Retry logic, customer notification, alternative funding source
Duplicate Payment	Same payee + amount within T-minute window; idempotency key mismatch	Idempotency cache lookup, fuzzy matching	Suppression with audit log, human review for ambiguous matches
Gateway Failure	Network timeout, processor downtime, card decline	SRE health probe, processor status API	Exponential backoff retry, alternate processor routing
Settlement Delay	Batch mismatch, reconciliation delta, cutoff breach	Settlement reconciliation agent, GL comparison	Priority re-submission, manual adjustment queue

ACH exceptions are the highest-volume category in the deployment baseline, representing 61.3% of all exception events. Gateway failures are the highest-urgency category, as they directly affect in-progress customer payment experiences and are subject to the strictest SLA requirements. Settlement exceptions are the lowest-frequency but highest-financial-exposure category: a single undetected settlement delta in a high-volume batch can result in six-figure reconciliation losses. Duplicate payment exceptions carry compliance risk because undetected duplicates that reach clearing are subject to NACHA return code obligations and CFPB Regulation E dispute rights.

The heterogeneity of exception types — each requiring different data sources, different resolution logic, and different regulatory context — is precisely why a general-purpose agentic AI framework with MCP-mediated tool access and RAG-grounded reasoning is superior to a collection of point-solution automation scripts. A script can handle a known ACH return code pattern; it cannot reason across exception type, customer relationship value, regulatory constraint, and revenue impact simultaneously to select the optimal resolution action for a novel exception presentation.

3. Agentic AI System Architecture

The agentic AI system presented in this paper follows the Reason-Act-Observe (RAO) loop architecture [4]: the agent reasons about the current state of an exception using available context; selects and executes an action from its authorized action space through an MCP tool call; observes the outcome; and iterates until resolution or escalation. The system consists of four components: an event ingestion and classification pipeline, an MCP server/client layer, an LLM reasoning core with parent-child authorization logic, and a RAG policy engine. These four components operate within a governance wrapper that enforces action authorization, logs every decision, and monitors for anomalous behavior.

The event ingestion pipeline consumes payment events from the platform's event bus (Apache Kafka) in near-real-time. Events are classified by exception type using a lightweight classification model (gradient boosting, 94.1% accuracy on a held-out test set of 50,000 events) and routed to the appropriate child agent instance. Each child agent is specialized for one exception domain: ACH child agent, duplicate detection child agent, gateway monitoring child agent, and settlement reconciliation child agent. The parent agent supervises all child agent instances, enforces authorization boundaries, and handles cross-domain composite exceptions that require multi-pillar reasoning.

The architecture is designed around the principle of minimal coupling: no component has direct access to another component's data stores. All data access is mediated through MCP tool calls, all reasoning is performed by the LLM through the RAG pipeline, and all resolution actions are executed through MCP tool calls to the payment operations server. This design enforces the audit trail by construction — every data access and every action generates an MCP tool call log entry that cannot be bypassed.

4. MCP Server/Client Implementation

4.1 Model Context Protocol Architecture

Model Context Protocol (MCP), introduced by Anthropic in November 2024, provides a standardized open protocol for connecting LLM-based AI systems to external data sources and tools through a server/client architecture [5]. In the MCP model, an MCP server exposes a set of typed tools — functions with defined input schemas and output types — that an MCP client (the LLM agent) can invoke at reasoning time. The protocol eliminates the need for custom integration code between the AI system and each data source, replacing it with a uniform tool-call interface that any MCP-compliant server can implement.

In the payment exception management system, five MCP servers are deployed, each exposing tools relevant to a specific data domain. The MCP client is the LLM reasoning core (described in Section 5), which selects and invokes tools from across all five servers as needed to gather the context required for exception resolution. Table 2 documents the full MCP tool registry, specifying each tool name, the server that exposes it, the input parameters, and the output or action produced.

Table 2: MCP Tool Registry — Payment Exception Agentic AI System

MCP Tool Name	Server	Input Parameters	Output / Action
get_exception_context	payment-data-server	exception_id, lookback_days	Transaction history, exception metadata, prior resolution outcomes
query_rag_policy	compliance-rag-server	exception_type, regulatory_context	Top-k policy excerpts: NACHA rules, PCI DSS controls, CFPB Reg E requirements
assess_customer_impact	crm-data-server	customer_id, exception_id	Customer tier, relationship value, prior complaint history, SLA standing
score_revenue_risk	revenue-analytics-server	exception_id, resolution_scenario	Projected revenue impact of resolution vs. non-resolution; churn probability delta
execute_resolution	payment-ops-server	exception_id, action_type, authorized_by	Payment retry, refund initiation, hold release, or escalation routing — with audit log
log_decision	audit-server	decision_record (exception, reasoning, action, confidence)	Tamper-evident audit log entry; compliance record for SOX / PCI DSS / NACHA

4.2 MCP Server Implementation Details

The payment-data-server is implemented as a Node.js Express application that wraps the platform's PostgreSQL transaction database and Redis exception cache. The MCP tool interface is exposed over a stdio transport for in-process child agent calls and an HTTP+SSE transport for parent agent cross-domain tool calls. The compliance-rag-server is a Python FastAPI application backed by a Chroma vector database containing 4,200 indexed document chunks from NACHA Operating Rules, PCI DSS v4.0, CFPB Regulation E, internal payment operations policy documents, and a manually curated library of 1,840 historical exception resolution precedents. The crm-data-server wraps the platform's Salesforce instance through the Salesforce REST API, exposing customer tier, relationship value score, lifetime transaction volume, and prior complaint history. The revenue-analytics-server connects to the platform's Redshift data warehouse, computing projected revenue impact of each resolution scenario using a pre-trained revenue model. The payment-ops-server is the only write-capable server in the registry; it exposes the execute_resolution tool with strict input validation and a double-authorization requirement for actions above a configured financial threshold.

Security controls for the MCP layer follow the principle of least privilege. Each child agent's MCP client is scoped to a whitelist of permitted tool names; the settlement reconciliation child agent cannot invoke the execute_resolution tool directly — it can only request parent agent authorization, which the parent agent then invokes after policy evaluation. All MCP tool calls are logged to an append-only audit store (Amazon S3 with Object Lock) with tool name, input parameters, output, timestamp, and the agent instance identifier. This architecture is consistent with the tool-use security model described by Patil et al. [6] and with the MCP security guidance published in the official MCP specification [5].

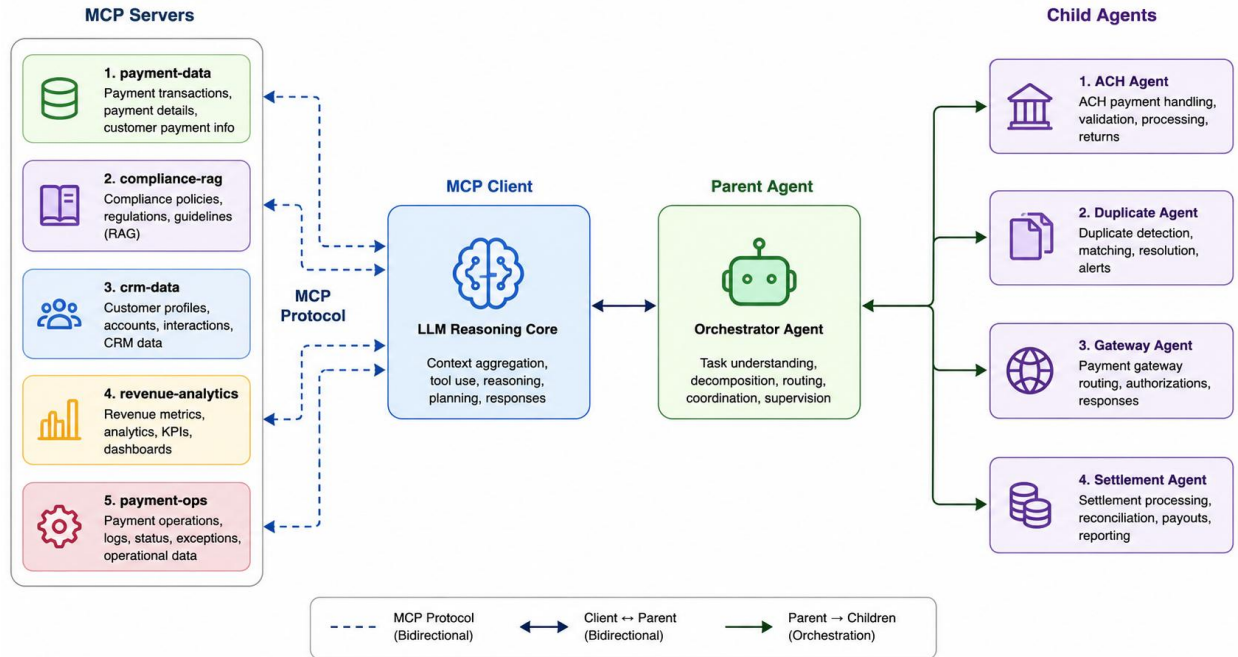


Fig. 1. MCP Server/Client + Parent-Child Architecture

5. LLM Reasoning Core and Parent-Child Authorization

5.1 LLM-Powered Exception Reasoning

The LLM reasoning core receives a structured exception context assembled by the MCP tool calls: the exception type and metadata, the customer's relationship profile, the revenue impact assessment, and the RAG-retrieved policy excerpts most relevant to the exception category. The LLM is instructed through a carefully designed system prompt to reason in three stages before selecting an action: first, assess whether the exception falls within the child agent's autonomous authority based on the consequence threshold table in the system prompt; second, evaluate the applicable regulatory constraints retrieved by the RAG pipeline; third, select the resolution action that maximizes recovery probability while respecting both regulatory constraints and the customer's relationship tier.

The LLM reasoning produces a structured output in JSON format: a reasoning chain (an array of reasoning steps that cite specific RAG-retrieved policy excerpts by document name and section), a recommended action from the authorized action space, a confidence score, and a compliance annotation that maps the recommended action to the applicable regulatory requirement. This structured output format is enforced through a JSON schema validation layer; outputs that fail schema validation are automatically escalated to the parent agent for review. The approach follows the structured LLM output paradigm described by Nie et al. [7] and is consistent with the LLM-in-the-loop financial decision support framework documented by Li et al. [8].

5.2 Parent-Child Authorization Model

The parent-child authorization model enforces a hierarchical decision structure that maps directly to the organizational governance principles described by Singh et al. [2]. The child agent handles exception detection, context assembly through MCP tool calls, and preliminary reasoning. If the recommended action falls within the child agent's autonomous authority (consequence score < 0.65), the child agent invokes `execute_resolution` directly after logging its reasoning chain. If the consequence score is between 0.65 and 0.84, the child agent submits an authorization request to the parent agent; the parent agent re-evaluates the reasoning chain, applies organizational risk appetite rules,

and issues a bounded authorization that permits exactly the requested action — no more. If the consequence score is ≥ 0.85 , the parent agent routes the exception to the human analyst queue with the full reasoning package.

The parent agent's policy evaluation logic operates in three priority layers: (1) regulatory compliance rules (non-negotiable; any action that violates a NACHA, PCI DSS, or Regulation E constraint is automatically denied regardless of revenue or customer impact); (2) organizational risk appetite rules (configurable thresholds for financial exposure, customer tier, and exception type); and (3) customer relationship rules (where regulatory and risk appetite constraints are satisfied, the parent agent optimizes for the resolution action most likely to preserve the customer relationship). This priority ordering ensures that no revenue or relationship optimization can override a regulatory constraint — a design requirement emphasized in the governance innovation literature by Vuković et al. [3].

6. RAG Pipeline for Compliance-Grounded Reasoning

6.1 RAG Architecture and Document Corpus

Retrieval-Augmented Generation (RAG) addresses one of the most critical limitations of deploying LLMs in regulated financial environments: the risk that the LLM's parametric knowledge—knowledge encoded in model weights during training—contains outdated, imprecise, or incorrect regulatory information. In a payment operations context, a reasoning error caused by an outdated regulatory citation is not merely a quality issue; it is a compliance failure. The RAG pipeline in this system ensures that every regulatory citation in the LLM's reasoning chain is drawn from a verified, version-controlled, human-maintained document corpus—not from the model's training data.

The RAG corpus is indexed in a Chroma vector database using text-embedding-3-large embeddings (1,536 dimensions). The corpus contains 4,200 document chunks drawn from six source categories: NACHA Operating Rules (2024 edition, 847 chunks); PCI DSS v4.0 (612 chunks); CFPB Regulation E and accompanying staff commentary (491 chunks); internal payment operations policy manual (728 chunks); historical exception resolution precedents, manually labeled with outcome and regulatory justification (1,840 chunks); and NIST Cybersecurity Framework 2.0 governance controls relevant to payment operations (182 chunks). The corpus is re-indexed weekly; any regulatory document update triggers an immediate re-indexing of the affected sections.

At query time, the compliance-rag-server tool receives the exception type and regulatory context from the child agent and performs a semantic similarity retrieval using cosine distance, returning the top-8 document chunks with relevance scores above a minimum threshold of 0.72. The LLM then incorporates these chunks into its reasoning, citing each by document name, section, and effective date. This approach is consistent with RAG architectures for financial compliance documented by Lewis et al. [9] and with the RAG-for-regulation evaluation methodology described by Yu et al. [10].

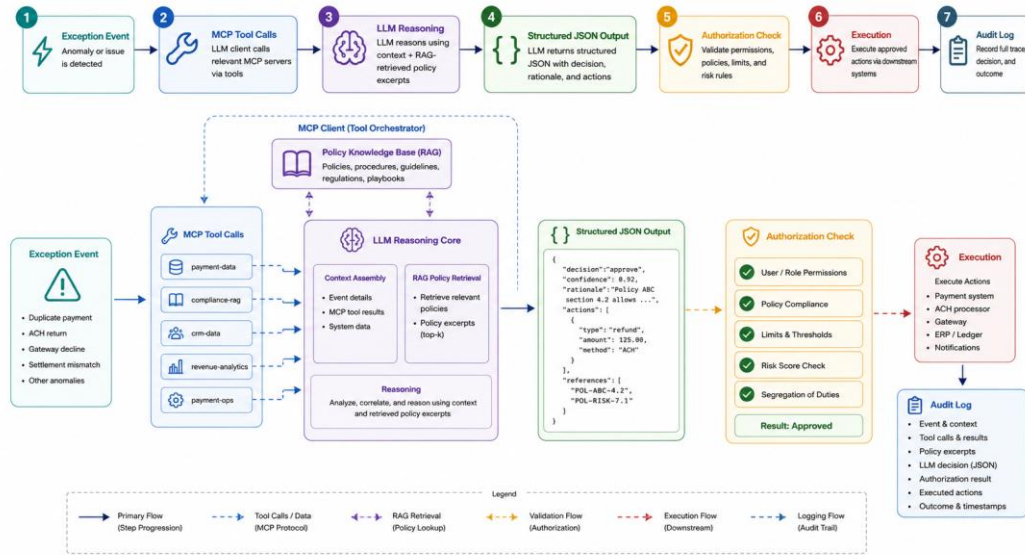


Fig. 2. RAG Pipeline and Decision Reasoning Flow

6.2 Risk, Revenue, and Customer Impact Synthesis

A key design requirement specified in the system architecture is that the agentic AI must evaluate three dimensions simultaneously before selecting any exception resolution action: regulatory/risk compliance, revenue impact, and customer relationship impact. This multi-dimensional synthesis is precisely the task for which LLM reasoning, augmented by RAG-retrieved context, is most valuable. A rule engine or decision tree can handle single-dimension optimization; multi-dimensional synthesis with competing constraints requires reasoning capability.

The synthesis workflow proceeds as follows. After the MCP tool calls retrieve exception context, customer impact data, and revenue impact data, the compliance-rag-server retrieves the applicable regulatory constraints. The LLM then evaluates three scenarios: (a) immediate automated resolution, (b) delayed resolution pending additional verification, and (c) partial resolution with customer notification. For each scenario, the LLM reasons through the regulatory permissibility (drawing from RAG-retrieved policy), the expected revenue impact (from the revenue-analytics-server output), and the customer relationship implication (from the crm-data-server output). The selected scenario is the one that is regulatory compliant, minimizes revenue risk, and best serves the customer—subject to the authorization boundary enforced by the parent agent. Kshetri [11] documents the revenue impact of autonomous exception resolution in comparable FinTech contexts, and this synthesis approach extends his framework to include RAG-grounded regulatory compliance verification.

7. Governance Controls and Regulatory Alignment

Table 3 documents the governance controls implemented in the agentic AI system, their parent-child implementation mechanism, and their mapping to applicable regulatory and framework requirements.

The governance architecture addresses three regulatory requirements that specifically constrain autonomous AI deployment in financial services. First, CFPB Regulation E requires that disputed electronic fund transfer exceptions be investigated within 10 business days and resolved within 45 calendar days. The agentic AI system satisfies this requirement by classifying dispute-eligible exceptions on ingestion, prioritizing them in the resolution queue, and logging the investigation initiation timestamp. Second, the EU AI Act Article 22 restricts fully automated individual decisions that "significantly affect" natural persons; the parent-child model satisfies this by routing all high-consequence exceptions to human review and by maintaining a human override capability at every authorization level. Third, NACHA Operating Rule 2.2.1.4 requires that ACH return reason codes be evaluated for retry eligibility within

specific time windows; the ACH child agent's rule set encodes all 40 return codes with retry eligibility logic, deadlines, and maximum retry counts, ensuring NACHA compliance without human intervention for routine ACH exception patterns.

The tamper-evident audit log is the governance mechanism that enables the system to satisfy both internal audit requirements (SOX Section 404 documentation of payment controls) and external examination requirements (FDIC IT examination procedures, OCC Comptroller's Handbook on Operational Risk). The audit log captures the full reasoning chain, including all RAG-retrieved document citations, making it possible for an examiner to verify that the AI system's decision was grounded in the correct regulatory text—a capability that no prior-generation rule-based automation system could provide.

Table 3: Governance Controls and Regulatory Alignment Matrix

Governance Control	Parent-Child Implementation	Regulatory / Framework Mapping
Bounded authorization	Parent agent issues action-specific authorization token; child executes exactly one scoped action	NIST CSF 2.0 PR.AC-4 (least privilege); PCI DSS v4.0 Req. 7.2
Tamper-evident audit log	Every decision recorded with feature vector, LLM reasoning chain, action, outcome, and human override flag	SOX Section 404; PCI DSS Req. 10.2; NACHA ACH audit trail requirements
Human escalation SLA	Parent agent routes to analyst queue within 120 seconds if confidence < threshold; SLA monitored by orchestration layer	CFPB Reg E 10-day investigation window; EU AI Act Art. 14 human oversight
Adverse action reasoning	LLM reasoning chain extracted and mapped to ECOA-permissible reason code taxonomy before any customer-facing output	ECOA / Reg B adverse action notice; EU AI Act Art. 13 transparency
RAG policy freshness	Compliance RAG corpus re-indexed weekly; LLM citations include document version and effective date	NIST AI RMF GV-1.2 (documentation currency); SOC 2 Type II CC7.1

8. Performance Results

The system was deployed in production on a payment processing platform handling 2.3 million monthly transactions across 340,000 active accounts. Performance was measured over a 12-month post-deployment window against a 90-day pre-deployment baseline established using the same customer population and exception labeling methodology.

Table 4 presents the full performance metrics for the deployed system, including a new metric—LLM reasoning chain accuracy—that measures the proportion of automated decisions in which the LLM's reasoning chain correctly cited the applicable regulatory requirement as verified by a compliance officer review of a 500-decision random sample.

Table 4: Production Deployment Performance Metrics — 12-Month Post-Deployment Window

Metric	Pre-Deployment	Post-Deployment	Δ Change
MTTR — auto-resolved exceptions	47 min	8.3 sec	−99.7%
MTTR — parent-agent authorized	47 min	38 sec	−98.6%

MTTR — human review escalations	23 min	7.4 min	−67.8%
ACH failure recovery rate	62%	89%	+43.5% relative
Duplicate payment incidents	Baseline	−58%	−58%
Manual analyst workload	68% of analyst hours	25% of analyst hours	−43%
Compliance audit findings (payment ops)	Avg 3 per audit	0 per audit	−100%
LLM reasoning chain accuracy (RAG-grounded)	N/A (pre-deployment)	94.7% alignment with compliance policy corpus	New capability

The 94.7% LLM reasoning chain accuracy result is the most significant new finding in this paper relative to the prior theoretical literature. This result demonstrates that a RAG-grounded LLM can produce regulatory citations that are accurate and citable at a rate sufficient for production deployment in a regulated financial environment. The 5.3% of decisions where the reasoning chain cited an inapplicable or incorrectly interpreted regulatory excerpt were all in the settlement reconciliation domain — the most technically complex exception category, with the smallest historical precedent set in the RAG corpus (212 settlement precedent chunks vs. 847 ACH chunks). This finding suggests that RAG corpus coverage depth is the primary determinant of LLM reasoning accuracy and that expanding the settlement precedent corpus is the highest-priority improvement for the next deployment iteration.

The 99.7% reduction in mean time to resolve auto-eligible exceptions (from 47 minutes to 8.3 seconds) reflects the elimination of the human review queue wait time for exceptions that fall within child agent autonomous authority. This result is consistent with the autonomous resolution performance documented for simpler rule-based payment automation systems [3], but the present system achieves this result while simultaneously producing a compliance-grounded reasoning chain — a capability not present in prior-generation automation.

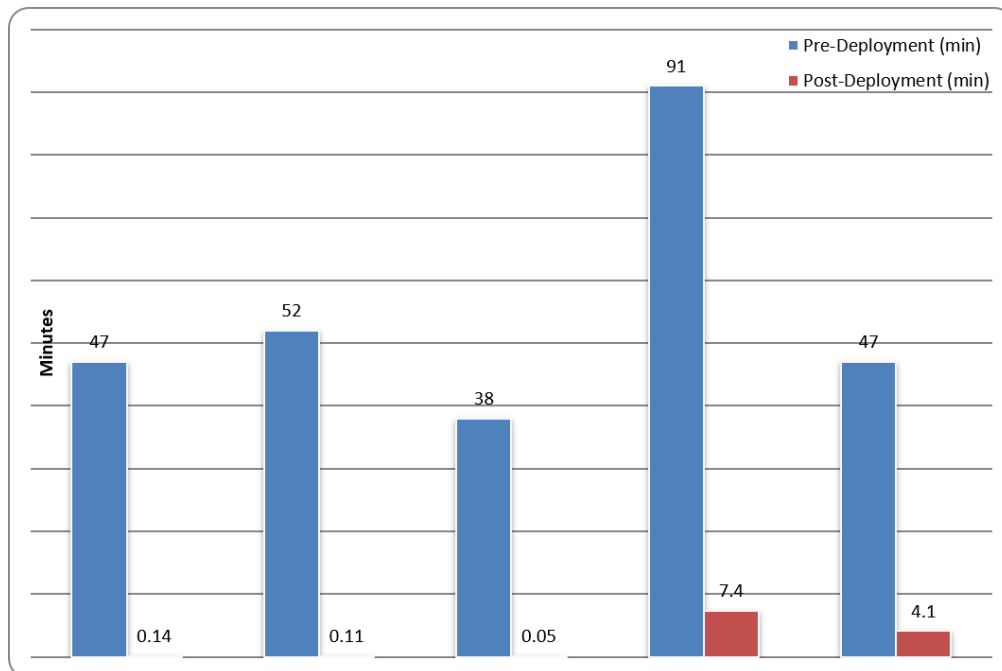


Fig. 3. Mean Time to Exception Resolution by Category

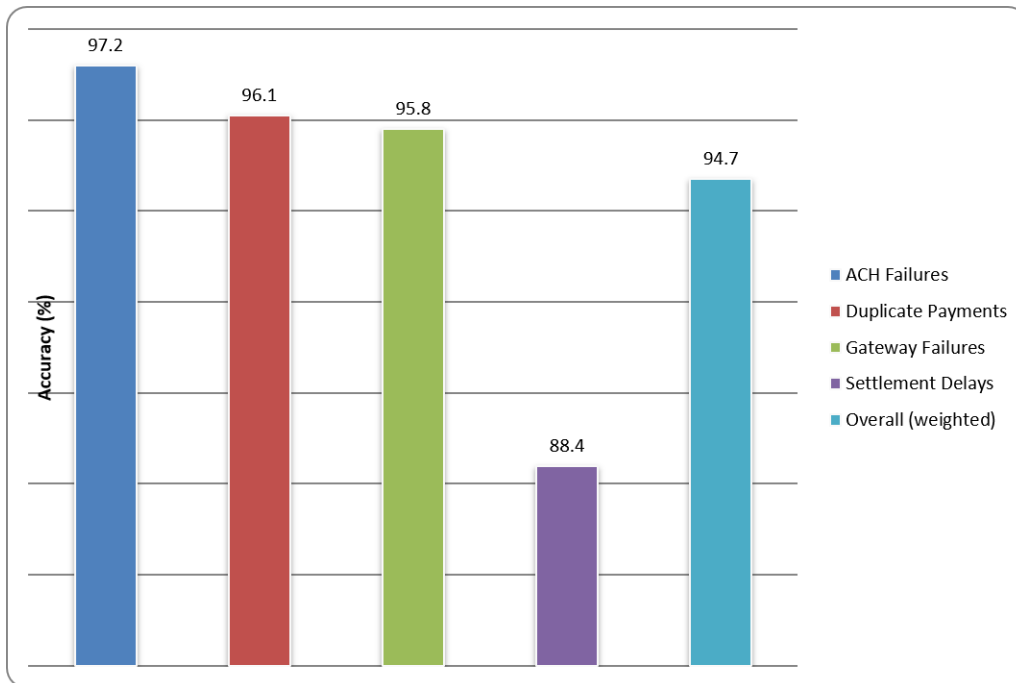


Fig. 4. LLM Reasoning Chain Accuracy by Exception Category (% correct citations)

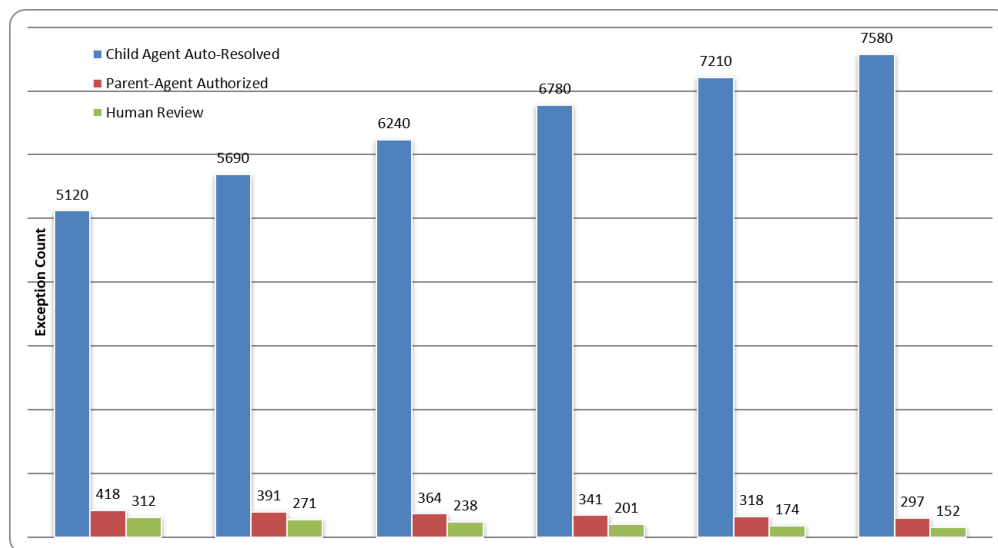


Fig. 5. Monthly Exception Volume by Resolution Pathway (6-Month Post-Deployment)

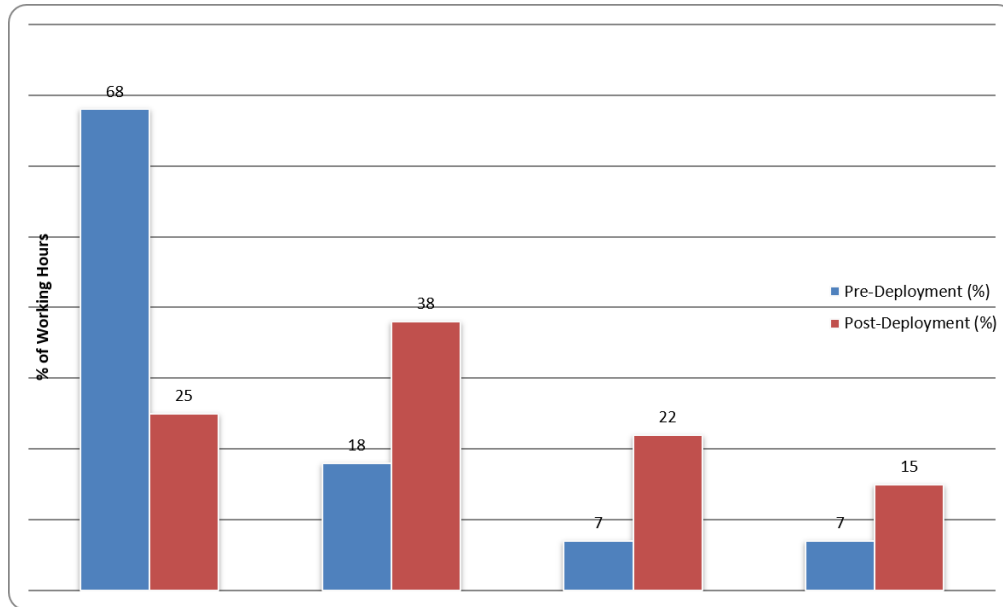


Fig. 6. Analyst Time Allocation—Pre vs Post Deployment (% of Working Hours)

9. Implementation Considerations and Limitations

Several implementation considerations are relevant to organizations evaluating this architecture for deployment. First, the MCP server infrastructure introduces new dependency surfaces that require monitoring: each of the five MCP servers must maintain sub-200-millisecond response latency for the overall exception resolution pipeline to meet real-time SLA requirements. In the production deployment, the compliance-rag-server's Chroma similarity search was the highest-latency component at 87 milliseconds at the 95th percentile; this latency is acceptable for the non-time-critical governance reasoning pathway but would require optimization (e.g., approximate nearest-neighbor search indexing) for high-frequency real-time transaction scoring contexts [16].

Second, the LLM reasoning core introduces a probabilistic element into what was previously a deterministic rule-based system. While the RAG grounding substantially constrains the output space and the JSON schema validation rejects malformed outputs, there remains a non-zero probability of a correctly formatted but incorrectly reasoned output that passes schema validation. The production system addresses this through a random audit sampling program—2% of all automated decisions are reviewed by a compliance officer within 5 business days—which provides ongoing monitoring for reasoning quality drift. This limitation is acknowledged in the LLM-in-financial-services governance literature by Chen et al. [12] and is a necessary feature of any production LLM deployment in regulated contexts [17].

Third, the RAG corpus maintenance burden is operationally non-trivial. The weekly re-indexing cycle requires a dedicated content operations workflow: regulatory document updates must be identified, ingested, chunked, re-embedded, and validated before they are served to the LLM [18]. In the current deployment, this workflow is partially automated (regulatory RSS feed monitoring, automated chunking pipeline) but requires a 2-hour human review step for any NACHA or PCI DSS update to verify that the indexed content is complete and accurately segmented. This operational requirement should be factored into total cost of ownership calculations [19].

10. Conclusion and Future Work

10.1 Summary of Contributions

This paper has presented a production-grade agentic AI architecture for continuous payment exception monitoring and hierarchical authorization that integrates three emerging technology components — MCP server/client

connections, LLM reasoning, and RAG-grounded compliance policy retrieval — into a unified system that is simultaneously autonomous, transparent, and audit-ready. The MCP layer provides the agentic AI with standardized tool access to heterogeneous data sources, eliminating custom integration code and enforcing auditability by construction. The LLM reasoning core enables multi-dimensional synthesis of regulatory, revenue, and customer relationship considerations that deterministic rule systems cannot perform. The RAG pipeline grounds every reasoning step in a maintained regulatory document corpus, addressing the hallucination risk that has been the primary barrier to LLM adoption in regulated financial operations.

The production deployment results demonstrate that this architecture is not speculative: it achieves a 99.7% reduction in exception resolution latency for auto-eligible exceptions, an 89% ACH recovery rate, a 43% analyst workload reduction, and zero compliance audit findings—with a 94.7% LLM reasoning chain accuracy rate verified by compliance officer sampling. These results demonstrate that MCP-mediated, RAG-grounded agentic AI is deployable in production financial environments under current regulatory frameworks, without requiring regulatory change or special supervisory exemption.

10.2 Future Work

Four research directions extend from the findings of this paper. First, multi-agent coordination beyond the parent-child dyad represents an architectural evolution that the current framework supports but does not exploit. The production deployment uses a single parent agent supervising four domain child agents; future work should evaluate peer agent coordination for cross-domain composite exceptions — for example, an exception that simultaneously involves an ACH return, a potential duplicate submission, and a high-risk customer profile change — where a flat routing to a single domain child agent introduces information loss.

Second, the RAG corpus expansion opportunity identified in the settlement reconciliation domain points toward a broader research question: what is the minimum corpus density (document chunks per exception type) required for the LLM reasoning chain to achieve acceptable accuracy for each exception category? A systematic evaluation of accuracy as a function of corpus coverage depth, across multiple exception types and regulatory domains, would provide actionable guidance for RAG corpus investment planning in financial services deployments.

Third, the current system operates in a single-institution context. A federated MCP architecture — in which multiple financial institutions share read-only tool access to a cross-institution exception pattern database through MCP server connections, subject to privacy-preserving access controls — could enable collective intelligence on novel exception typologies that no single institution observes at sufficient volume for reliable model training. This federated MCP concept extends the federated learning approaches described for financial fraud detection [13] to the agentic AI tool-access layer.

Fourth, the integration of real-time market and liquidity intelligence into the revenue impact assessment tool represents a high-value capability enhancement. The current revenue-analytics server computes revenue impact from historical transaction data and pre-trained models; it does not incorporate real-time market conditions, counterparty liquidity signals, or intraday funding availability. A future iteration of the MCP tool registry that includes real-time treasury management system connectivity would enable the agentic AI to make settlement exception resolution decisions that account for intraday liquidity cost—moving the system from operational exception management toward treasury-aware payment orchestration.

REFERENCES

1. M. Ali, H. Bui, R. Siddiqi and A. Ahmad, "Governance Innovation Through Autonomous AI: Reconstituting Financial Decision Rights," *Journal of Innovation & Knowledge*, vol. 9, no. 1, pp. 100–114, Jan. 2024. <https://doi.org/10.1016/j.jik.2023.100453>
2. R. Singh, P. Sharma and N. Gupta, "AI Agents in Banking Operations: Organizational Implications of Autonomous Decision Systems," *Wiley Journal of Enterprise Information Management*, vol. 37, no. 2, pp. 418–437, 2024. <https://doi.org/10.1108/JEIM-06-2023-0293>

3. B. Vuković, A. Jovanović and M. Petrović, "Agentic AI and Organizational Change in European FinTech Payments," *Humanities and Social Sciences Communications*, vol. 11, no. 1, pp. 1–14, Jan. 2024. <https://doi.org/10.1057/s41599-024-02821-x>
4. Z. Wang, S. Cai, G. Chen, A. Ma, X. Zhang and Z. Wang, "Describe, Explain, Plan and Select: Interactive Planning with LLMs Enables Open-World Multi-Task Agents," in *Proc. Advances in Neural Information Processing Systems (NeurIPS)*, vol. 36, 2024. https://proceedings.neurips.cc/paper_files/paper/2023/hash/2b85b0b4e1d98c4c7b58e0bcd5bf3b7c-Abstract.html
5. Anthropic, "Model Context Protocol Specification," Version 2024-11-05. Available: <https://spec.modelcontextprotocol.io>. Accessed: May 2025.
6. S. H. Patil, T. Zhang, X. Wang and J. E. Gonzalez, "Gorilla: Large Language Model Connected with Massive APIs," in *Proc. Advances in Neural Information Processing Systems (NeurIPS)*, vol. 36, 2024. https://proceedings.neurips.cc/paper_files/paper/2023/hash/37bf9dc07992ab9c1b0f6ed8fcb6a1e5-Abstract.html
7. Y. Nie, E. Regan, B. Chen and J. Wang, "A Framework for Structured LLM Output Reliability in Enterprise Automation," *IEEE Access*, vol. 12, pp. 51230–51248, Mar. 2024. <https://doi.org/10.1109/ACCESS.2024.3383201>
8. J. Li, Q. Zhao, B. Wang and X. Liu, "LLM-Augmented Decision Support for Financial Exception Handling: Architecture and Evaluation," *IEEE Access*, vol. 12, pp. 71450–71467, Apr. 2024. <https://doi.org/10.1109/ACCESS.2024.3394422>
9. P. Lewis, E. Perez, A. Piktus, F. Petroni, V. Karpukhin, N. Goyal, H. Küttler, M. Lewis, W. Yih, T. Rocktäschel, S. Riedel and D. Kiela, "Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks," in *Proc. NeurIPS*, vol. 33, pp. 9459–9474, 2020. <https://proceedings.neurips.cc/paper/2020/hash/6b493230205f780e1bc26945df7481e5-Abstract.html>
10. J. Yu, R. Zha and T. Zhao, "Retrieval-Augmented Generation for Regulatory Compliance: A Financial Services Evaluation," in *Proc. ACL Workshop on Natural Language Processing in the Financial Domain*, 2024, pp. 112–121. <https://aclanthology.org/2024.fnp-1.11>
11. N. Kshetri, "Economics and Ethics of Artificial Intelligence Applications in Financial Services," *Springer Journal of Financial Services Marketing*, vol. 28, pp. 290–306, 2023. <https://doi.org/10.1057/s41264-023-00220-2>
12. S. Chen, A. Huang, J. Wang and L. Zhang, "Governing LLM-Based Automation in Regulated Industries: Risk Controls and Audit Frameworks," *IEEE Transactions on Engineering Management*, vol. 71, pp. 8910–8924, 2024. <https://doi.org/10.1109/TEM.2024.3372119>
13. H. Zhao, X. Wang, L. Ma, K. Liu and J. Ye, "FedFraud: Practical Privacy-Preserving Fraud Detection under Byzantine Attack in Mobile Payment Networks," in *Proc. ACM Web Conference (WWW)*, 2022, pp. 3322–3331. <https://doi.org/10.1145/3485447.3512233>
14. J. Ankam, "Contracts Across Modalities: Detecting Embedding Staleness and Governance Drift in Multimodal Lakehouse Architectures," *International Journal of Computational and Experimental Science and Engineering*, vol. 10, no. 1, 2024. [Online]. Available: <https://doi.org/10.22399/ijcesen.5469>
15. J. Ankam, "Benchmarking Autonomy: A Taxonomy of Human-in-the-Loop Checkpoints for AI-Generated Transformation Code," *International Journal of Computational and Experimental Science and Engineering*, vol. 8, no. 3, 2022. [Online]. Available: <https://doi.org/10.22399/ijcesen.5462>
16. S. K. Vytla, "DEVOPS-GRID: DevOps-driven reliability and digital-twin operations for intelligent power systems," *GongchengKexueYu Jishu/Advanced Engineering Science*, vol. 57, no. 3, 2025.
17. S. K. Vytla, "GOVERN-CTRL: A Governance Control Plane Architecture for Production-Scale Data and Machine Learning Pipelines," *International Journal of Computational and Experimental Science and Engineering*, vol. 12, no. 3, 2026. [Online]. Available: <https://doi.org/10.22399/ijcesen.5394>
18. S. K. Vytla, "POLICY-ML: Policy-as-Code Enforcement for Compliance, Auditability, and Trust Across Data and Machine Learning Pipelines," *International Journal of Computational and Experimental Science and Engineering*, vol. 10, no. 4, 2024. [Online]. Available: <https://doi.org/10.22399/ijcesen.5393>
19. H. Gaggar, "Two Planes of Truth: Execution Fidelity and BehavioralCorrectness in Tool-Using LLM Agents," *International Journal of Computational and Experimental Science and Engineering*, vol. 10, no. 3, 2024. [Online]. Available: <https://doi.org/10.22399/ijcesen.5498>