

FINANCIAL SERVICES · AGENTIC AI · DATA INTELLIGENCE

Conversational BI at Scale: A Financial Services Blueprint for Governed Multi-Agent AI

*Unifying Fragmented Data Ecosystems with
Agentic AI for the Modern Financial Enterprise*



AUTHORS

Shivkumar Krishnan, Biplab Adhikary, Sathiyarayanan
Altimetrik · AI Engineering

PREPARED BY

Altimetrik · AI Engineering

CONTENTS

Inside this whitepaper

01	Executive Summary	03
02	Business Impact & The Problem	04
03	What's Broken Today	05
04	The Hidden Gaps & The Cost of Inaction	06
05	The Turning Point	08
06	Altimetrik's Approach & Four Core Capabilities	09
07	The Framework	11
08	Technical Challenges & How We Solved Them	12
09	Platform Architecture & Maturity Model	15
10	A View of How Success Looks & Real-World Scenario	16
11	How to Get Started, Checklist & The Bottom Line	18
12	About Altimetrik	21

01 OVERVIEW

Executive Summary

Enterprise data platforms in financial services sit at the intersection of regulatory scrutiny, operational risk, and analytical demand. Yet the data ecosystems that power credit ratings, market intelligence, indices, and infrastructure operations remain fragmented, with critical metadata, asset inventories, and KPIs scattered across discovery tools, asset-management systems, BI platforms, and bespoke spreadsheets.

In a global financial information and analytics enterprise, this fragmentation manifests across more than 50,000 physical servers and dozens of data domains. Multiple discovery sources (a Configuration Management Database (CMDB), an endpoint discovery tool, and a cloud security posture tool) each provide a different view of the same estate. Metadata definitions drift across tools and there is no unified semantic layer that downstream systems can trust.

Altimetrik partnered with this enterprise to build an AI-driven, multi-agent platform that transforms fragmented data into governed, conversational intelligence. Anchored in Databricks, Microsoft Fabric, and a custom multi-agent framework, the solution delivers a unified server-estate gold dataset, automated metadata generation with confidence-scored certification, a Supervisor Agent for dynamic Genie orchestration, and an RDF-based knowledge graph for downstream semantic consumption.

By replacing fragmented manual workflows with intelligent orchestration, the platform turns data visibility, governance, and conversational analytics into a competitive advantage.

By replacing fragmented manual workflows with intelligent orchestration, the platform turns data visibility, governance, and conversational analytics into a *competitive advantage*.



02 THE CHALLENGE

Business Impact & The Problem

50K+

physical servers unified across 3
discovery sources

~70%

reduction in manual
reconciliation and metadata
effort

Weeks → Mins

for infrastructure reconciliation cycles

The Problem: When Data Velocity Outpaces Governance

Financial services enterprises run on data. Credit ratings depend on it. Market intelligence is built on it. Financial indices are derived from it. And the infrastructure that powers all of this, tens of thousands of servers, hundreds of pipelines, dozens of analytical workspaces, must remain visible, secure, and governed in real time.

Over the last decade, the volume and velocity of this data have accelerated. Cloud-native platforms, hybrid infrastructure, vulnerability scanners, FinOps tools, and conversational AI have each added value, but each tool ships with its own data model, its own refresh cycle, and its own definition of the truth.

While the analytical surface has transformed, the underlying data fabric has not.



03 THE CHALLENGE

What's Broken Today

In most large financial enterprises, the operating model still relies on fragmented, manual processes to keep data trustworthy. A simple question ("How many servers do we have? What's our vulnerability exposure? Which dataset feeds this dashboard?") triggers a chain of human-dependent steps. Engineers reconcile counts across tools. Analysts hand-build lineage. Subject-matter experts certify metadata in long review cycles. This creates systemic friction at every stage:

- 01 **Inconsistent inventories:** Multiple discovery and asset-management platforms (a CMDB, an endpoint discovery tool, and a cloud security posture tool) produce different counts for the same server estate, with limited reconciliation capability.
- 02 **Fragmented metadata:** Schema, definitions, lineage, and ownership are stored in silos across Databricks, Power BI, the CMDB platform, and security tools, with no shared definitions or refresh cycles.
- 03 **Siloed AI agents:** Conversational BI agents (Genies) operate independently, require manual provisioning, and have no central intelligence layer to interpret user intent or route requests.
- 04 **No semantic backbone:** Without a unified ontology, cross-table relationship discovery remains manual, and downstream systems cannot consume governed semantic artifacts.

As ecosystems grow more interconnected, even a small gap in visibility can ripple across compliance reporting, vulnerability response, FinOps optimization, and analytical trust. Without a reliable way to unify, govern, and converse with this data, teams either move slowly or take risks they cannot fully control.



04 THE CHALLENGE

The Hidden Gaps

When a business question actually enters the system ("What's our infrastructure footprint by region? Which servers are out of compliance? What does this KPI mean and where does it come from?") It rarely moves with the speed the business expects.

Instead, it slows down.

A request lands in a ticket or a Teams channel. A data engineer or platform analyst must first interpret what the question really means. Context is scattered across CMDB exports, vulnerability dashboards, FinOps reports, and tribal knowledge. The definitions are unclear. What looks like a simple inventory question may require reconciliation across three discovery tools, plus FinOps cost data, plus vulnerability posture, each with its own refresh lag.

Teams begin to proceed cautiously.

Reconciliation becomes a manual exercise: time-consuming, and often incomplete. Metadata certification depends on individual SME availability rather than standardized workflow. Conversational BI agents work, but only within the narrow scope of a single Genie. And analytical artifacts that downstream systems need (knowledge graphs, ontologies, certified datasets) either don't exist or exist only in spreadsheets.

And so, what should take hours takes days. **What should take days takes weeks.**

Many financial enterprises still treat this as a backend hygiene problem, a matter of better dashboards, more headcount, or one more reconciliation script. But the real challenge is structural: the data fabric was never designed for the speed, scale, and conversational expectations of modern financial operations. Excessive caution leads to delayed audit readiness, reactive vulnerability response, and a growing trust gap between the business and its own data.

WHERE THIS LEAVES

Financial Enterprises

The industry has reached a tipping point. Financial information leaders are operating in ecosystems where infrastructure footprints span tens of thousands of servers and multiple discovery sources, regulatory and audit demands are continuously increasing, and the cost of inconsistency is immediate - for compliance, for security, and for trust.

Yet, the underlying approach to managing data visibility and governance remains rooted in manual, fragmented processes designed for a very different era. This creates a structural gap, one that cannot be solved through incremental improvements.



This creates a structural gap, one that cannot be solved through incremental improvements.

THE COST OF

Inaction

The consequences of this broken model are no longer just technical. They are business critical.

Inconsistent infrastructure inventories delay audit responses, slow vulnerability remediation, and undermine FinOps optimization. In a regulated industry, this directly translates to compliance risk and unbudgeted spending.

More critically, fragmented metadata is a leading cause of analytical mistrust. When KPI definitions drift between tools and lineage cannot be traced, dashboards become unreliable, AI agents return inconsistent answers, and business users lose confidence in their own data.

The financial cost is severe. Manual reconciliation absorbs senior engineering time. Audit cycles stretch. Conversational AI investments stall because agents cannot scale beyond manual provisioning.

As data velocity continues to increase, this gap between business expectation and operational capability will only widen. What was once a manageable inefficiency has become a structural bottleneck, limiting scalability, increasing risk, and holding back AI adoption.

The Turning Point

To keep up with the pace of modern financial operations, enterprises must move beyond viewing data unification, metadata governance, and conversational AI as separate backend functions. They must converge into an intelligent, scalable, and reliable orchestration layer, one that can match business speed without compromising control.

Because in today's financial services landscape, the ability to converse with your data, accurately, in context, and at scale, is no longer an operational concern.

It is a competitive advantage.

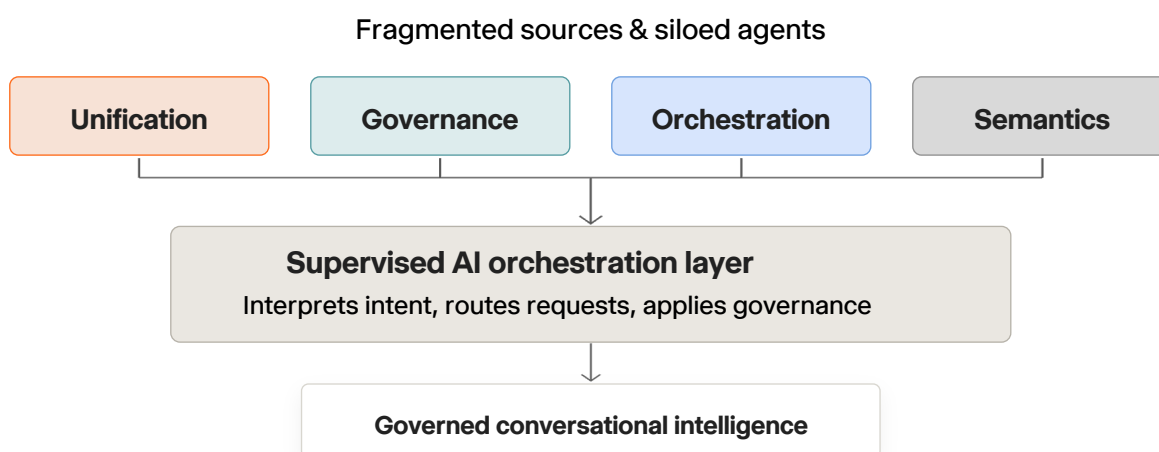
06 THE SOLUTION

Altimetrik's Approach: An AI-Orchestrated Framework for Data Intelligence

The problem with fragmented data ecosystems isn't just inefficiency. It is that the systems themselves were never designed for the scale, speed, and conversational nature of modern financial analytics.

Altimetrik's approach is not an incremental improvement. It is a fundamental rethinking of how data is unified, governed, and consumed, shifting from manual reconciliation to intelligent orchestration.

At the core of this approach is a supervised, AI-driven framework that combines four critical capabilities:



06 THE SOLUTION

Four Core Capabilities

01

Unification

UNDERSTANDING THE ESTATE

A unified, automated CI Server Estate Gold Dataset reconciles a CMDB, an endpoint discovery tool, and a cloud security posture tool into a single governed view, exposing reconciled inventory, vulnerability posture, and FinOps cost in Power BI dashboards and Microsoft Fabric.

02

Governance

UNDERSTANDING THE MEANING

Automated metadata generation backed by a confidence-score certification workflow, an LLM evaluation and regeneration loop, and notification-based approvals, producing a metadata repository that is the foundation for the Enterprise Knowledge Graph.

03

Orchestration

UNDERSTANDING THE INTENT

A Supervisor Agent provides a central intelligence layer that interprets user intent, routes requests to the right Genie, dynamically creates new Genies on demand from metadata, and manages the lifecycle and disposal of temporary agents.

04

Semantics

UNDERSTANDING THE RELATIONSHIPS

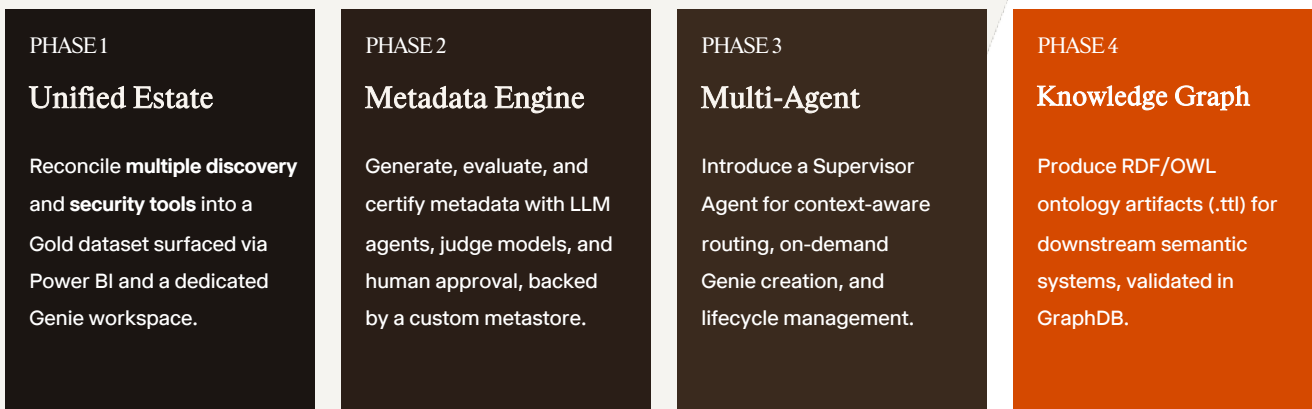
An RDF-based Knowledge Graph models cross-dataset relationships in an ontology-driven format, automating synonym resolution and producing structured .ttl artifacts that downstream systems can ingest and validate via GraphDB.

07 THE SOLUTION

The Framework:

The platform is delivered as a multi-phase program that converts fragmented data sources into a governed conversational and semantic layer through an automated, supervised workflow.

You can think of this as a 4-phase intelligent pipeline:



FRAGMENTED SOURCES

GOVERNED INTELLIGENCE

This transforms data unification, governance, and conversational AI from fragmented, human-driven processes into a structured, repeatable, and scalable system.

08 ENGINEERING

Technical Challenges and How We Solved Them

Building a unified, governed, conversational data platform at this scale surfaced several non-trivial engineering problems. The following challenges were the ones that most shaped the architecture, and the approaches taken to resolve them illustrate the deeper technical character of the platform.

1 Reconciling fragmented infrastructure inventory across three discovery tools

THE CHALLENGE

A CMDB, an endpoint discovery tool, and a cloud security posture tool each produce a different view of the same server estate. Schemas differ, server identifiers do not align, refresh cycles are inconsistent, and asset definitions vary by tool. Naive joins produced both duplicates and gaps. Manual reconciliation in spreadsheets was the only reliable path, and it could not keep up with the rate of change.

OUR APPROACH

We built a Bronze, Silver, and Gold medallion pipeline on Delta Lake. The Bronze layer ingests raw source data with full lineage. The Silver layer applies source-specific transformations and harmonizes server identifiers using deterministic matching rules combined with custom resolution logic. Custom PySpark jobs then compute reconciled KPIs in the Gold layer, governed by Unity Catalog and exposed to Power BI through Microsoft Fabric OneLake shortcuts. The Gold layer refreshes continuously, replacing the prior manual cycle with a governed, lineage-tracked view.

2 Trusting AI-generated metadata at enterprise scale

THE CHALLENGE

Auto-generating metadata across hundreds of tables and thousands of columns using LLMs introduces real hallucination risk. Inconsistent column descriptions, fabricated business glossary terms, or incorrect lineage claims would erode trust faster than the manual processes they replace. Subject-matter experts cannot realistically review every generated artifact at this scale.

OUR APPROACH

We designed a multi-stage evaluation pipeline leveraging Claude and Open AI LLMs. A primary Claude Sonet LLM agent generates metadata (descriptions, business glossary terms, lineage hints, tags). A separate OpenAI Judge LLM evaluates each artifact independently, scoring confidence on dimensions such as completeness, terminology consistency, and alignment with the source schema. Low-confidence outputs trigger automatic regeneration with refined prompts. High-confidence outputs route to SMEs through SMTP notification for review in the NextGen UI. From there, SMEs have two paths: approve the metadata as-is, in which case it persists into the custom metastore, or provide additional context such as clarifications, corrections, or missing business nuance, and request regeneration. In the regeneration path, the workflow re-runs with the SME's inputs incorporated into the generation prompt, producing a refined version for another review cycle. Only approved metadata persists alongside Unity Catalog and is exposed to downstream consumers via REST API.

3 Resolving attribute naming inconsistency across entities

THE CHALLENGE

Similar attributes appeared under inconsistent names across entities, and many lacked any documented meaning at all. A single business concept like server cost might surface as `srv_cost`, `cost_amt`, `monthly_charge`, or `total_spend` across different tables, with no shared definition tying them together. Business users could not find or interpret attributes consistently, which slowed every analytical question and undermined trust in self-service BI. Even with metadata generated, the inconsistent naming meant users still struggled to know which attribute meant what across the estate.

OUR APPROACH

We built a canonical naming layer driven by metadata-based synonym discovery. The system profiles each column using both its metadata (descriptions, comments, tags) and a sample of the underlying data, then uses an LLM-assisted process to identify columns that represent the same underlying business concept across different entities. Candidate synonyms are clustered, and a canonical name is derived for each concept and persisted in the custom metastore. Business users see a single canonical attribute name regardless of which underlying table they query, and AI agents (including Genies) use the canonical layer to resolve attribute references consistently across the platform.

4 Scaling conversational AI beyond native platform limits

THE CHALLENGE

Databricks Agent Bricks Supervisor enforces a hard limit of 20 sub-agents per supervisor and 30 tables per genie space. With dozens of business divisions, data domains, and analytical use cases, even modest growth would exceed this ceiling. Native Agent Bricks also does not support dynamic agent creation, which means every new Genie requires manual provisioning, registration, and lifecycle management.

OUR APPROACH

We built a custom Supervisor Agent on LangGraph that orchestrates Genies through a registry pattern rather than a fan-out subscription. User queries, route to the Supervisor, and trigger a vector search against a Genie Registry to identify candidate agents. A Judge LLM breaks ties when multiple agents could plausibly answer. When no suitable agent exists, the Supervisor consults custom metadata to identify the right tables and columns, then provisions a new on-demand Genie scoped to that data. Agent lifecycle (registration, usage tracking, disposal of temporary agents) is fully automated, and MLFlow provides observability across the orchestration layer.

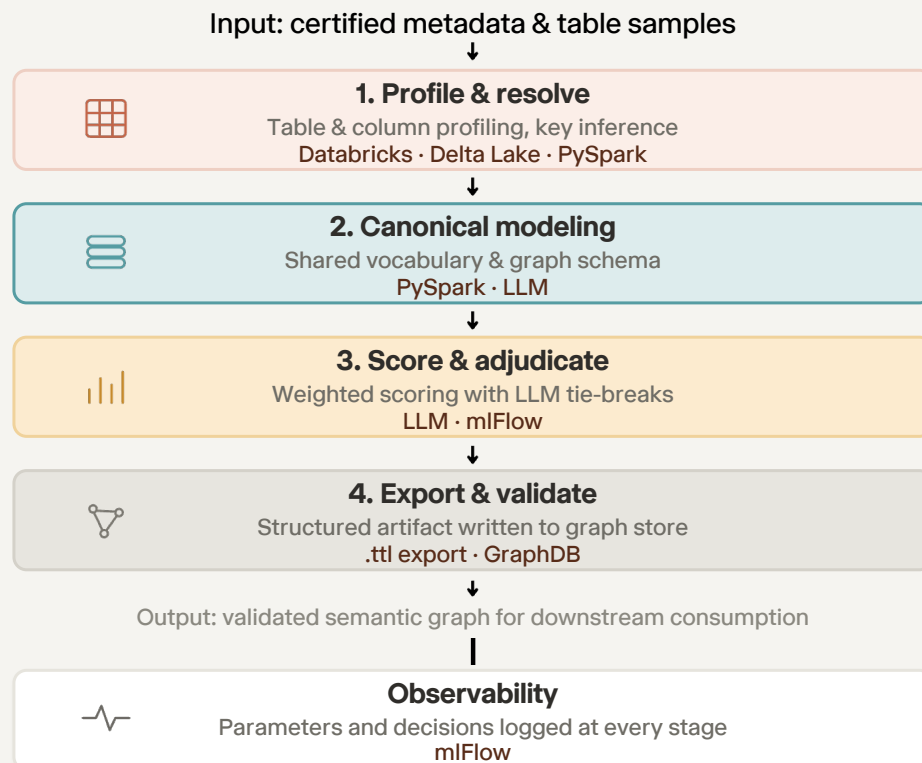
5 Generating governed semantic artifacts for downstream consumption

THE CHALLENGE

Downstream systems require RDF/OWL ontology artifacts to consume metadata in a structured semantic format. Producing these manually is slow and inconsistent. No existing pipeline could convert the platform's tabular metadata into a validated semantic graph that downstream systems could ingest reliably.

OUR APPROACH

We built a multi-stage RDF pipeline. Table resolution and column profiling characterize each dataset. Primary-key and foreign-key inference establish relationships. The canonical naming layer described in the previous challenge provides a consistent attribute of vocabulary, while contract building captures the structural agreements between sources. A graph schema is then constructed using a weighted scoring model across the inferred relationships, with an LLM tie-breaker for ambiguous cases. The pipeline logs parameters at every stage for observability and exports a structured .ttl artifact to predefined cloud storage. The artifact is then ingested and independently validated in GraphDB.



Knowledge-graph generation pipeline

Figure 2 — Pipeline stages from metadata profiling through export to a validated graph store, with observability logged throughout.

09 ENGINEERING

Platform Architecture & Maturity Model

Platform Architecture

The platform implements a supervised automation loop where AI agents perform the heavy lifting of reconciliation, generation, and orchestration, while humans retain control over certification, approval, and execution. This design balances automation benefits with the governance requirements of a regulated financial environment.

The architecture consists of three primary layers:

PRESENTATION LAYER

React-based NextGen UI plus Microsoft Teams and Power BI for submitting requests, reviewing reports, and approving certified metadata.

ORCHESTRATION LAYER

LangChain and LangGraph-powered agent pipelines plus a custom Supervisor Agent driving multi-agent workflows on Databricks (with MLFlow for observability).

DATA LAYER

Bronze, Silver, and Gold Delta Lake medallion architecture, Unity Catalog governance, a custom metastore, Microsoft Fabric OneLake shortcuts, and an RDF graph store (GraphDB).

These layers communicate through well-defined APIs and Genie endpoints, enabling modularity and independent scaling of components.

Maturity Model: From Manual to Autonomous

The framework provides a clear evolution path:

- | | |
|-----------|---|
| Level 1-2 | Manual reconciliation and SME-driven metadata certification. |
| Level-3 | AI-assisted unification, automated metadata generation, and supervised Genie orchestration. |
| Level 4-5 | On-demand agent creation, continuous knowledge graph optimization, and semi-autonomous semantic governance. |

This allows organizations to adopt AI incrementally, without disrupting existing governance models.

10 OUTCOMES

A View of How Success Looks

The true value of this approach is not theoretical; it is measurable and repeatable. Across the engagement, the platform demonstrates a step-change in performance:

01

Speed

Server-estate reconciliation cycles that traditionally took weeks of manual cross-tool comparison are reduced to continuous, refreshable Gold dataset views, with conversational answers available in Genie within seconds.

02

Quality

Automated metadata generation paired with confidence-scored certification and an LLM evaluation and regeneration loop reduces ambiguity in KPI and column definitions, eliminating the silent inconsistencies that erode analytical trust.

03

Efficiency

DBA, SME, and platform engineering workload is reduced by an estimated 70 percent, shifting their role from manual reconciliation and certification to oversight, governance, and strategic optimization.



10 OUTCOMES

Real-World Scenario: A Global Financial Information and Analytics Leader

In a publicly traded global financial information and analytics enterprise, with business divisions spanning Credit Ratings, Market Intelligence, Technology and Infrastructure BI, the data ecosystem supports tens of thousands of servers and dozens of downstream systems.

BEFORE

- Server inventory was distributed across a CMDB, an endpoint discovery tool, and a cloud security posture tool with inconsistent counts and limited reconciliation.
- Metadata across Databricks, Power BI, and the discovery and security tools lived in silos with no shared definitions or refresh cycles.
- Conversational BI agents (Genies) operated independently, requiring manual provisioning and offering no central intent routing
- No unified semantic model existed to support downstream systems requiring ontology artifacts

WITH THE ALTIMETRIK PLATFORM

- A unified CI Server Estate Gold Dataset reconciles all three discovery sources with continuous refresh, exposed via Power BI and a Genie conversational workspace
- Metadata is generated automatically by LLM agents, evaluated by a Judge LLM, certified through a confidence-score workflow, and stored in a custom metastore
- A Supervisor Agent (built on LangGraph) orchestrates Genie agents, dynamically creates on-demand Genies from metadata, and manages their lifecycle, overcoming the 20-subagent hard limit of the native Agent Bricks Supervisor
- An RDF-based Knowledge Graph produces .ttl artifacts for downstream consumption, validated in GraphDB.

THE BUSINESS IMPACT

Unified visibility across servers and 3 discovery sources

Higher consistency and trust in KPIs across business divisions

Conversational AI scaled beyond native platform limits

In essence, fragmented data ecosystems evolve from a bottleneck into a business accelerator.

11 ROADMAP

How to Get Started: Checklist & The Bottom Line

Adopting AI-driven data unification and multi-agent intelligence is not a one-step transformation. It is a structured journey. Altimetrik recommends a phased approach to ensure both impact and adoption.

PHASE 1 Pilot and Validate

- Identify a low-risk domain (for example, infrastructure inventory)
- Deploy the unified Gold dataset for a limited scope
- Validate accuracy of reconciliation, metadata generation, and Genie responses
- Build internal confidence with measurable outcomes

PHASE 2 Expand and Standardize

- Extend to multiple business divisions and data domains
- Enrich the metadata repository with business-approved definitions
- Build and maintain the Enterprise Knowledge Graph
- Introduce governance checkpoints and SME approval workflows

PHASE 3 Integrate and Scale

- Deploy the Supervisor Agent with on-demand Genie creation
- Enable self-service conversational BI for analysts and business users
- Integrate with CI/CD and data engineering workflows

PHASE 4 Optimize and Evolve

- Publish RDF/OWL artifacts to downstream semantic systems
- Use historical agent interactions to improve recommendations
- Move toward continuous validation and semi-autonomous metadata governance

11 ROADMAP

Checklist for Getting Started

- ✓ Assess current data unification, metadata, and conversational AI maturity
- ✓ Identify high-friction, high-impact domains (infrastructure, FinOps, vulnerability)
- ✓ Build foundational assets (custom metastore, knowledge graph schema)
- ✓ Start small with a single phase, measure impact, and scale progressively
- ✓ Align teams on AI governance, SME certification, and approval workflows

The Bottom Line

Financial enterprises don't need to replace their current data ecosystems. They need to *augment them with intelligence*.

By adopting Altimetrik's AI-driven, multi-agent framework, financial information and analytics leaders can:

- Match the speed and conversational expectations of modern financial operations
- Reduce risk and inconsistency without slowing down innovation
- Unlock the full value of conversational AI investments at enterprise scale

— About AI-Powered Financial Data Platforms

This whitepaper describes a delivered engagement combining AI-assisted data unification, automated metadata governance, multi-agent orchestration, and semantic knowledge modeling for enterprise financial data platforms. The work demonstrates the feasibility of combining modern AI techniques with the Databricks lakehouse, Microsoft Fabric, and ontology-driven semantic systems to deliver measurable business value.

For more information, including technical deep dives, implementation guides, and case studies, please contact aifirst@altimetrik.com.

ACKNOWLEDGMENTS

This work reflects the contributions of ~20-member delivery team of engineers, data scientists, architects, and program managers who designed, built, and operationalized the platform. The authors gratefully acknowledge Dipak Kumar Ghosh (Client Partner) for sponsoring the engagement and building strong customer partnership; Nitesh Jain (Product Owner) for translating business needs into the structured requirements that shaped the platform; Pramita Athikary (Program Manager) for the program leadership that held the delivery together across phases; and Rizwan Patel, Pranabesh Sarkar & Sandeep Kalra, for their foundational contributions to the Data and AI architecture in the early phases of the initiative. The authors also thank the broader delivery team whose engineering, design, and operational excellence made the outcomes described in this paper possible.

DISCLAIMER

This whitepaper describes an engagement-specific implementation and should not be construed as a turnkey production solution. Organizations considering similar implementations should conduct thorough security reviews, performance testing, and risk assessments appropriate to their specific environments and requirements. Quantitative outcomes referenced (for example, percentage reductions and scale figures) are derived from the engagement and may vary in other contexts.



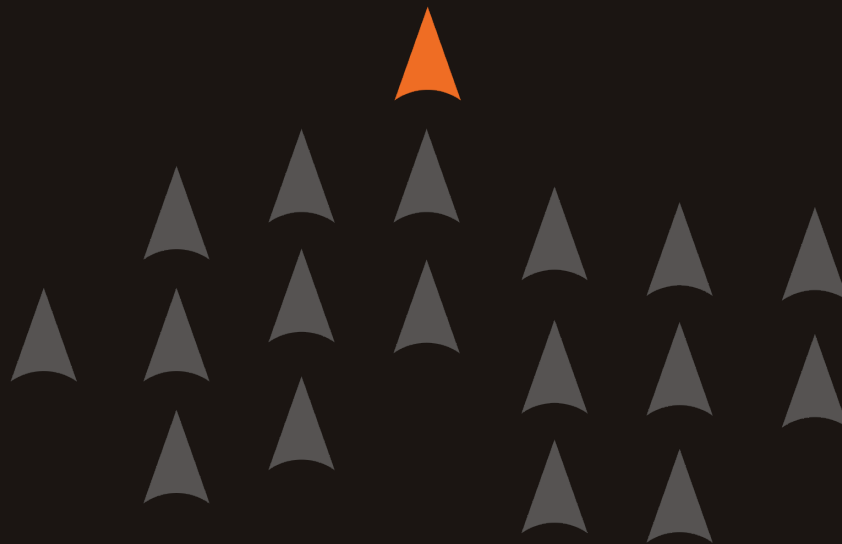
About Altimetrik

Altimetrik is an AI engineering company, building the systems that power the modern enterprise.

We are Altimetrik, an AI-native engineering company helping some of the most revered and iconic enterprises modernize systems, data, and processes at the heart of their business, so they can move faster, operate more efficiently, and innovate continuously. Through ALTi AIOS, our AI-native operating system, we combine the latest AI capabilities with deep engineering expertise to help clients solve complex challenges, accelerate modernization, and deliver measurable outcomes at scale.

Our clients get access to the latest AI innovations while maintaining the flexibility to choose the right technologies for their business, through trusted relationships with OpenAI, Google Gemini, Anthropic, Databricks, AWS, and Azure.

A member of the World Economic Forum's Centre for AI Excellence, the Forum's global hub for shaping responsible AI, Altimetrik is also recognized in the 2025 Constellation Research ShortList for Global AI Services and named a Major Contender in multiple Everest Group PEAK Matrix® assessments, including Software Product Engineering Services (2026), Enterprise Quality Engineering Services (2025), and Digital Engineering Services for BFSI and Life Sciences. Learn more at altimetrik.com.



Connect with us:

aifirst@altimetrik.com

www.altimetrik.com