

Fiddler Enterprise Agentic Observability

See every action. Understand every decision. Control every outcome.



IAS

Thumbtack

LENDINGPOINTAMERICAN FAMILY
INSURANCE

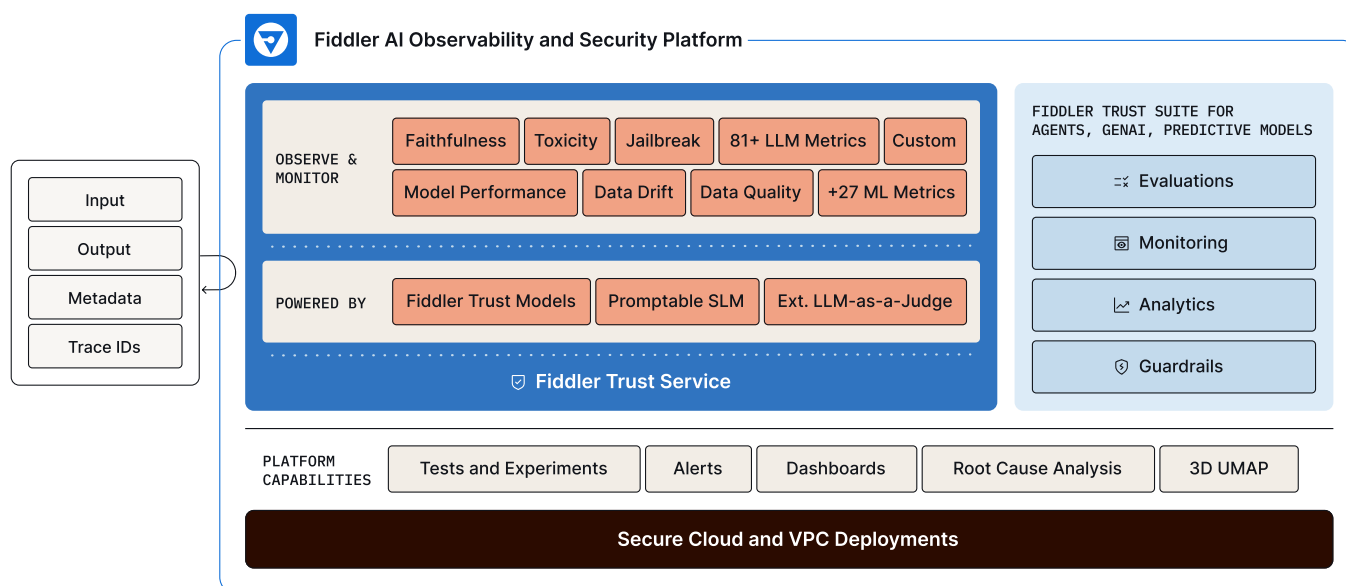
Fiddler delivers enterprise-grade, end-to-end observability for agentic AI. It enables organizations to build reliable, cost-effective agents with a continuous feedback loop across the lifecycle from evaluation and monitoring to analytics.

It gives AI and engineering teams visibility, context, and control across autonomous and reflective multi-agent systems. Teams can see every layer of the agentic hierarchy (application, session, agent, trace, and span) with aggregate rollups and granular insights on reasoning and decision paths, tool calls, and token usage, without sifting through millions of logs.

Fiddler Trust Service: The Foundation of Agentic Observability

☒ NATIVE, NOT ADDITIVE☒ NO EXTERNAL CALLS☒ ENTERPRISE SECURE

Fiddler Agentic Observability is powered by the Fiddler Trust Service, with purpose-built, secure, cost-effective Trust Models that run entirely in your environment. These Trust Models score agent and LLM inputs and outputs, enabling real-time monitoring without the need for external API calls, so there are no add-on fees or hidden costs.



Complete Visibility, Context, and Control



Deliver High Performance AI

Build, test, monitor, and improve to deliver contextually-aware agentic systems. Gain insights from system-wide health to span-level metrics.



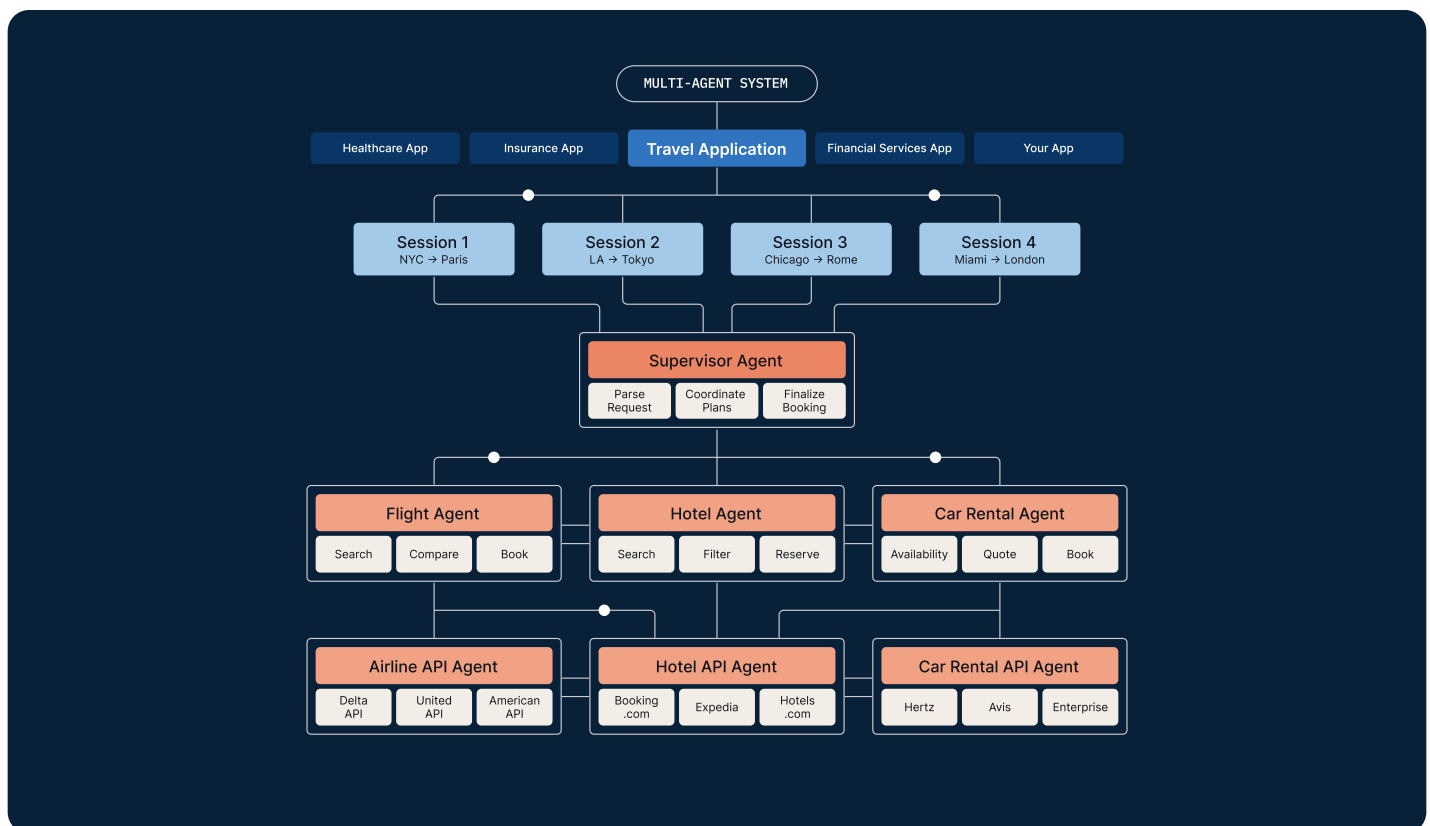
Protect From Costly Risks

Ensure safe operations at scale by detecting harmful exposure. Prevent costly brand and reputational risks.



Maximize ROI

Measure accuracy and cost side-by-side by testing and tuning models and configurations to achieve the best performance per token.



Key Enterprise Capabilities for the Agentic Observability Lifecycle

Name	Description	Application	Dataset	Created By	Created At	Status	Metadata
eval_demo-2023-10-05-01	Comprehensive evaluation of LLM GBA...	lm_app_evaluation	qa_evaluation_data...	admin@hubble.ai	October 5, 2023 11:...	Completed	model_type: mistral, 30k; evaluator_prompt: 5; evaluation_version: v1.0
eval_demo-2023-10-05-02	Comprehensive evaluation of LLM GBA...	lm_app_evaluation	qa_evaluation_data...	admin@hubble.ai	October 5, 2023 11:...	Completed	model_type: mistral, 30k; evaluator_prompt: 5; evaluation_version: v1.0
eval_demo-2023-10-05-03	Comprehensive evaluation of LLM GBA...	lm_app_evaluation	qa_evaluation_data...	admin@hubble.ai	October 5, 2023 12:...	Completed	model_type: mistral, 30k; evaluator_prompt: 5; evaluation_version: v1.0
hubble_ai_eval-2023-10-05-04	Comprehensive evaluation of chatbot GBA...	lm_app	hubble_eval_2	admin@hubble.ai	October 5, 2023 12:...	Completed	model_type: mistral, 30k; evaluator_prompt: 7; evaluation_version: v1.0
reg_eval_demo-2023-10-05-05	Comprehensive LLM evaluation	reg_eval_demo	reg_eval_demo	admin@hubble.ai	October 5, 2023 12:...	Completed	-
eval_demo-2023-10-05-06	Comprehensive evaluation of chatbot GBA...	lm_app_evaluation	qa_evaluation_data...	admin@hubble.ai	October 5, 2023 12:...	Failed	model_type: mistral, 30k; evaluator_prompt: 5; evaluation_version: v1.0
eval_demo-2023-10-05-07	Comprehensive evaluation of chatbot GBA...	lm_app_evaluation	qa_evaluation_data...	admin@hubble.ai	October 5, 2023 12:...	Failed	model_type: mistral, 30k; evaluator_prompt: 5; evaluation_version: v1.0

Evaluation

Build cost-effective, reliable agents by evaluating in pre-production to improve performance early and reduce post-launch risks.

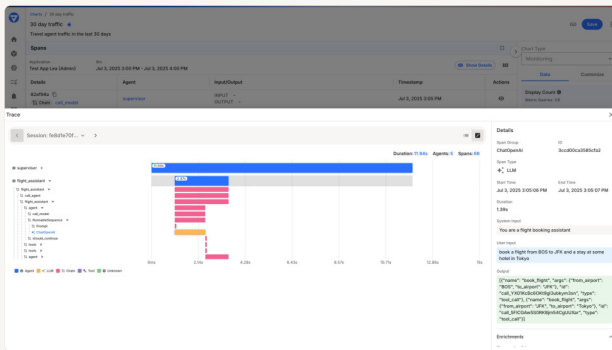
- Evaluate agents with golden and challenger datasets to improve performance and behavior.
- Run experiments with different prompts and responses and compare side-by-side to reach better outcomes.
- Stress-test edge-cases to surface weaknesses and points of failure modes early.



Monitoring

Gain real-time monitoring in a single-pane command center.

- Get visibility across the agentic hierarchy: application → session → agent → trace → span to see what happens in any interaction.
- See aggregate-to-granular insights in customizable dashboards and reports.
- Monitor key metrics including hallucination, toxicity, PII/PHI, jailbreak, and custom KPIs.
- Trace reasoning chains, tool calls, and decision paths across sessions.



Analytics

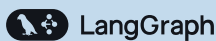
Understand why issues happen, then close the loop from production to development to continuously improve agents.

- Diagnose issues at every level of the agentic hierarchy all the way to pinpointing the failing span, reducing MTTI and MTTR.
- Use filters, sorting, and span attributes to expose bottlenecks and critical issues across tools, chains, agents, and LLMs.
- Perform root cause analysis to identify errors and visualize the exact breakdown of issues that affected agent performance.

OTel-Compatible and Framework-Agnostic

Fiddler integrates natively with OpenTelemetry (OTel) to preserve your existing AI stack. OTel unifies data and custom attributes across systems so Fiddler can deliver deep visibility into agent performance and behavior without changing your pipeline.

Compatible with:



Custom Agentic Framework

Fiddler is the all-in-one AI Observability and Security platform for responsible AI. Our evaluations, monitoring and analytics capabilities provide visibility, context and control across development and production. This gives teams actionable insights to build better, more reliable AI agents, and GenAI and ML applications. An integral part of the platform, the Fiddler Trust Service provides quality and moderation controls for AI agents and GenAI applications. Powered by cost-effective, task-specific, and scalable Fiddler-developed Trust Models — they deliver the fastest guardrails in the industry. Fiddler offers flexibility in secure deployment options through cloud and VPC environments.

Fortune 500 organizations use Fiddler to scale AI agents, GenAI, and ML deployments. This helps them deliver high performance AI, avoid costly AI risks, and maximize ROI.