



Artificial Intelligence Finance Institute

AIFI · PROFESSIONAL EDUCATION

PROFESSIONAL CERTIFICATE · OCTOBER–NOVEMBER 2026 · 2026 CURRICULUM

AI Agents in Finance Certificate

Twenty-one core instructional hours, plus seven optional 60-minute office hours — 28 live hours available. Build six lab systems and one evaluated capstone — seven working systems in total, with the evidence a risk function needs to assess a controlled pilot.

Taught by **Miquel Noguer i Alonso, PhD, ARPM** — Founder & Chief Science Officer, AIFI
· with **David Pacheco Aznar**, Staq.io

DURATION

**28 Hours
Available**

21 core + 7 optional office hours

FORMAT

Live Online

lab-driven, hands-on

LEVEL

Advanced

practitioner-level

DATES

Oct 12–Nov 23

Mondays · 2026

Seven working systems, not seven demos

This is a build course. Each of the seven sessions opens with a short concept block and spends the rest of its three hours in the editor: you leave every session with an agent that runs, a test that proves it runs, a trace that shows what it did, and a note on what it cost. The arc goes from a single tool-using assistant to a multi-agent investment committee that survives being killed mid-decision, then through observability, retrieval, memory and reinforcement-learning layers to a benchmarked prop-desk crew, and finally the evaluation and governance work that decides whether any of it is fit for a controlled pilot.

Labs use production frameworks — OpenAI Responses SDK, AutoGen, CrewAI, LlamaIndex, LangGraph, Letta, Claude Agent SDK, Google ADK and Hermes Agent — alongside the agentic reinforcement-learning stack (ToolRL, ReTool, ARPO, verl) and simulators (TradingAgents, FinRL-Meta). Hermes Agent — the MIT-licensed open agent framework from Nous Research, distinct from the Hermes model family the same lab publishes — is the reference harness, and it is taught across three declared styles rather than one: **stateless-parser mode** (empty toolset whitelist, no persistent memory, strict JSON output at every proposal boundary) as the reproducibility and audit baseline, **skills-augmented**, and **persistent memory with self-verification and completion contracts**. What each style costs in dispersion, spend and auditability is measured in session one rather than asserted. The deterministic core, audit trail and risk gate stay framework-independent.

NON-NEGOTIABLE STANDARDS ENFORCED THROUGHOUT

Compute, don't state

Every derived number is computed in code; every source-reported number stays attached to its source.

Every action a span

Tool calls, sources, policy checks, exceptions and hand-backs are emitted as OpenTelemetry GenAI spans from session one — redacted before write, portable across backends, reused by every later lab.

Stoppable — and resumable

Budget caps, kill-switches and human hand-back points are wired in from the first session, joined by checkpoints, idempotent side effects and a rehearsed resume-from-crash drill. Stopping is easy; restarting safely is the exam.

THE AIFI HARNESS THEORY — THE SPINE OF THE COURSE

An agent is a kernel plus a harness, and the harness is the part you can actually specify

Raw kernel vs agent kernel

One model call under a fixed decoding contract, versus model, planner, tool router, tools, memory, parser, retries and runtime controls as a framework instantiates them. Two objects; two different things to test.

Substrate neutrality

The same declared agent on Hermes Agent, LangGraph or AutoGen is the same agent only relative to a declared environment class and history law. A framework migration is a substrate change, like a model swap.

What a harness must supply

Toolset whitelisting, schema-enforced output you can reject, a pinnable content-addressed configuration, native audit hooks — and telemetry you can export in a standard shape. Past that it is engineering convenience, not capability.

Two questions at each boundary

Does the agent kernel beat the raw kernel on return and on dispersion? Is a nominally identical harness operationally coherent across base models, harness styles and framework versions?

A second, operational spine runs underneath it. Session one stands up the OpenTelemetry pipeline every later lab emits into; session two the durability substrate every later lab resumes on. By session seven the capstone's monitoring dashboard is reading six weeks of the cohort's own traces — the evidence pack assembles itself as the course runs.

SESSION RHYTHM

Seven weekly live sessions on Mondays, each three hours: a short concept block, the lab build, and a debrief where crews compare results. Seven optional 60-minute office hours follow immediately on the same Monday, for debugging and code review with the instructor who led that session; office hours are additional to the 21 core instructional hours.

TOOLS & SETUP

A current Python environment, API access to at least one frontier model provider, and a vector store. A starter repository — datasets, notebooks, agent scaffolds, an OTLP collector, a checkpoint store, microVM images and the evaluation harness — is issued before session one, with a setup check you run in advance.

The work that only happens when someone has time

A desk's binding constraint is not analysis, it is attention. Monitoring stops when the analyst is in a meeting. The second-order check gets skipped because the first one took the morning. What changes here is not that the work gets cleverer. It is that work which used to depend on someone being free now runs whether or not anyone is.

IF YOU ARE A ...	WHAT ONLY HAPPENS WHEN SOMEONE HAS TIME	WHAT RUNS WHETHER OR NOT ANYONE DOES
Analyst or researcher	The filings you got to before the call	Every filing in the coverage universe read on arrival, with the exceptions surfaced rather than buried
Portfolio manager	A committee memo written by whoever was free	An analyst, a sceptic and a chair producing the memo, with the dissent recorded rather than averaged away
Credit or risk	Covenant checks at quarter end	Covenants monitored as each 10-Q lands, breaches flagged with the page cited
Model validation	A review of the model you were handed	A traced decision record for every run — tool calls, memory writes and skill changes — with redaction and human hand-back already tested against attack
Desk engineer	A tool a colleague has to be shown how to use	A permissioned MCP endpoint any agent on the desk can call under a whitelist, inside an isolation tier you chose on purpose

WHAT SEPARATES A DEMO FROM SOMETHING A DESK CAN KEEP

Level 1	The agent produces an answer	
Level 2	The answer is reproducible	← most demos stop here
Level 3	The cost and latency of producing it are known	
Level 4	It survives a poisoned document, an adversarial prompt and a poisoned memory	
Level 5	A second-line reviewer can read the trail and sign off	

Levels three to five are where agent projects die inside institutions — not because the idea failed, but because nobody built the evidence while building the agent. Seven times over, this course builds both at once. What this edition changes is where that evidence comes from: spans instead of logs, checkpoints instead of restarts, scheduled memory instead of an append-only transcript. Levels three to five become properties the system exhibits continuously, not artefacts assembled the week before review.

WHO SHOULD ATTEND

Quantitative researchers, traders, risk and model-validation staff, and engineering leads in asset management, banking and fintech who need agent systems that survive a review, not a prompt-engineering tour.

- ✓ Buy-side and sell-side analysts and quants
- ✓ Risk, model-validation and compliance technologists
- ✓ Engineers building agentic features on a desk

PREREQUISITES

- Fluent Python — you write and debug your own code and are comfortable with APIs and async calls — plus working knowledge of probability, statistics and core finance (portfolio mathematics, time value of money, basic derivatives), and familiarity with machine-learning concepts including how a language model is trained and prompted.
- ✓ Advanced in finance and Python; no prior agent-framework experience assumed
 - ✓ No prior OpenTelemetry, durable-execution or microVM experience assumed — each is taught in the session that uses it
 - ✓ Prior reinforcement learning is helpful but not required
 - ✓ AIFI's Advanced Claude and GPT for Finance intensive is a natural precursor, but is not required

STRUCTURE

Schedule at a glance

DAY	DATE	NEW YORK	EUROPE	SESSION	HRS	YOU LEAVE WITH
Mon	Oct 12	09:00–12:00	15:00–18:00 CEST	1. Foundations: workflows, agents & harness styles	3.0	Agent scaffold + cost budget on OTel telemetry
Mon	Oct 19	09:00–12:00	15:00–18:00 CEST	2. Multi-agent architectures & durable orchestration	3.0	Investment-committee crew, crash-tested
Mon	Oct 26	09:00–12:00	14:00–17:00 CET	3. Observability, interpretability & decision tracing	3.0	Traced, redacted advice agent + memory audit
Mon	Nov 2	09:00–12:00	15:00–18:00 CET	4. Tool-using & coding agents: sandboxes, MCP & skills	3.0	Repair loop + MCP desk tool, in a microVM
Mon	Nov 9	09:00–12:00	15:00–18:00 CET	5. Context engineering & advanced memory (RAG 2.0)	3.0	Covenant monitor + memory-scheduling policy
Mon	Nov 16	09:00–12:00	15:00–18:00 CET	6. LLM agents + RL layers in financial simulation	3.0	Benchmarked prop-desk crew + skill-distillation arm
Mon	Nov 23	09:00–12:00	15:00–18:00 CET	7. Evaluation, governance & controlled pilots	3.0	Capstone + audit pack + live monitoring and ops runbook
Core instruction — plus seven optional 60-minute office hours					21.0	Six lab systems + capstone

Each session is followed by an optional 60-minute office hour, 12:00–13:00 New York — 18:00–19:00 CEST on 12 and 19 October; 17:00–18:00 CET on 26 October; and 18:00–19:00 CET from 2 November onward. Office hours are additional to the 21 core instructional hours and are held by the instructor who taught that session.

Europe's clocks go back on 25 October and New York's on 1 November, which is why the European column moves twice while the New York column does not. Faculty split: Noguera i Alonso leads sessions 1, 2, 5 and 7; Pacheco Aznar leads sessions 3, 4 and 6.

Taking both October courses? Advanced Claude and GPT for Finance runs 6–15 October. Its sessions 4 and 5 fall on Tuesday 13 and Thursday 15 October, in the same week as session one of this certificate — three sessions in five days on three separate days. Plan the workload for that week accordingly.

Session details

1 Foundations: workflows, agents & harness styles

Deciding what should be an agent, what should stay a script — and how much the harness is allowed to remember.

MON 12 OCT · 3 HOURS · NOGUER I ALONSO

You ship: a running agent scaffold, an OTel-instrumented cost dashboard, and a harness-style decision rule for your desk

TOPICS	The three tiers — deterministic workflow, AI-augmented workflow, autonomous agent — and the cost, latency and risk each buys. Raw kernel versus agent kernel, and why the harness is the object you can specify, pin and test. The harness-style spectrum on Hermes Agent: stateless-parser mode (empty whitelist, no persistent memory, strict JSON — the reproducibility and audit baseline) → skills-augmented → persistent memory with self-verification, completion contracts and a background curator. Mixture-of-Agents as a pickable ensemble, distinct from a crew. Token and cost budgeting before a line of orchestration is written. Choosing among Hermes Agent, OpenAI Responses SDK, AutoGen, CrewAI, LlamaIndex, LangGraph, Claude Agent SDK and Google ADK on grounds other than fashion. Instrumentation is OpenTelemetry-native from hour one: agent, workflow, tool and model spans, gen_ai.* token and latency metrics, convention version pinned.
LAB	Stand up the course repository and a research assistant that pulls market data and drafts a one-page note — then run the same declared agent under three harness styles and measure output dispersion, cost and auditability across them. The OTLP pipeline built here is reused by every later lab.
DELIVERABLE	A working scaffold, a cost dashboard reading standard spans, and a workflow-versus-agent and style-selection decision rule written for your own desk.
ASSIGNMENT	Graded lab 1 — the scaffold, the cost dashboard and the decision rule, reproducible from a clean clone.

2 Multi-agent architectures & durable orchestration

When a crew beats a single agent, when it just multiplies the bill — and how it survives a crash mid-decision.

MON 19 OCT · 3 HOURS · NOGUER I ALONSO

You ship: a three-agent investment committee that produces a defensible memo — and provably resumes from a kill without duplicating a side effect

TOPICS	Crew roles, delegation and typed hand-offs so agents cannot corrupt each other's state; consensus collapse, runaway delegation and cost blow-up as the failure modes. Production economics: orchestrator-worker as the default topology; roughly 58% token overhead for independent crews and near 285% for centralised ones; debate at about 2.5× single-model cost. A2A for cross-framework delegation; background subagent fan-out. Debate protocols and swarm topologies are surveyed rather than built, to make room for durable execution as the orchestration substrate: journal replay versus database checkpointing; framework-level checkpoints against engine-level durability; the replay-safety contract; idempotency keys on every side-effecting tool call; Saga-style compensation; a human approval as a durable signal rather than a blocked thread.
LAB	Build the analyst-sceptic-chair committee over SEC filings, dissent recorded rather than averaged away. Then the kill test: terminate the chair mid-memo, resume from checkpoint, and prove from the trace that no tool call re-fired and no decision was re-issued.
DELIVERABLE	The committee crew, its memo template, a cost-and-quality comparison against the single-agent baseline, and a durability posture note: what survives a crash, what does not, and the evidence.
ASSIGNMENT	Graded lab 2 — the crew, a three-configuration cost comparison, and a passed kill test with no side effect re-fired.

3 Observability, interpretability & decision tracing

Making an agent's decisions legible to someone who has to review them — including what it has learned about you.

MON 26 OCT · 3 HOURS · PACHECO AZNAR

You ship: an advisory agent whose every decision — and every memory write — leaves a readable, standards-based trail

TOPICS	Structured decision records capture tool calls, source traces, policy checks and exceptions without exposing hidden reasoning — mapped onto the OpenTelemetry GenAI conventions, so the decision record is the trace. Redaction enforced inside the instrumentation wrapper before any span is written, so client data and MNPI never reach the telemetry backend. Cost attribution per decision. The trace-to-eval loop: promoting production traces into version-controlled golden datasets a second line can replay. Memory governance: journey-style audits of accumulated memories and skills; memory and skill poisoning as first-class threats beside prompt injection; when a reviewer can trust the agent's own completion evidence. Pausable workflows and explicit human hand-back points.
LAB	Instrument the advice agent to the gen_ai.* conventions; verify its traces, policy checks and exception log; test that redaction and hand-back survive an adversarial prompt at the telemetry layer, not just the chat layer; then poison the memory store and confirm the audit surfaces it.
DELIVERABLE	A traced, redacted advisory agent, a decision-record schema mapped to the GenAI conventions, and a memory-audit procedure for your institution.
ASSIGNMENT	Graded lab 3 — the traced agent, its decision-record schema, and a memory-poisoning audit that surfaces the write.

4 Tool-using & coding agents: sandboxes, MCP & skills

Agents that write, run and repair code without escaping the sandbox — on the isolation ladder and the stateless protocol.

MON 2 NOV · 3 HOURS · PACHECO AZNAR

You ship: a self-repairing code agent in a microVM, one MCP server on the 2026 spec, and one audited skill

TOPICS	Coder–Critic–Executor loops (Reflexion, Self-Refine) and where each one stalls; unit-test-driven repair as the stopping criterion; stateless–parser boundaries at every proposal hand-off. The isolation ladder as a risk decision: hardened containers for reviewed code → gVisor for syscall interception → Firecracker and Kata microVMs for untrusted execution; snapshot-restore as the sandbox-layer mirror of workflow checkpointing; egress control, resource limits, defence in depth. MCP on the 2026-07-28 specification: the stateless core, the Tasks extension for long-running work, MCP Apps, OAuth/OIDC-aligned authorization, and trace-context propagation that lands tool calls on the same telemetry spine as everything else. Code execution with MCP: tools as code, tool search and progressive disclosure, filtering data before it reaches the model. The Agent Skills standard: SKILL.md anatomy, skills as procedural knowledge over MCP's connectivity — and as audited dependencies with code-execution rights.
LAB	A repair bot iteratively fixes a faulty VaR calculator against a failing test suite inside a microVM with explicit egress rules; you then wrap a redacted copy of one of your own pricing or reporting functions as an MCP server on the stateless spec under permission control, author one skill and review a peer's as a dependency audit.
DELIVERABLE	A working repair loop, an MCP server exposing a real desk tool safely, a sandbox policy (isolation tier, egress, limits) written for review, and one audited skill.
ASSIGNMENT	Graded lab 4 — the repair loop in its microVM, the MCP server on the 2026 spec, one skill and one peer audit.

5 Context engineering & advanced memory (RAG 2.0)

Grounding agents in filings and internal documents without inventing citations — and deciding, by priority, what an agent gets to remember.

MON 9 NOV · 3 HOURS · NOGUER I ALONSO

You ship: a streaming covenant monitor with page-resolving citations, an explicit context budget, and a priority-scheduled memory policy

TOPICS	Context engineering, within the episode: context rot; the context window as a capped per-turn budget; select / compress / order / isolate assembly; compaction with provenance-preserving summaries; state offloading with retrievable handles; sub-agent context isolation; prompt-caching economics and cache-hit telemetry. Advanced memory, across episodes: hierarchical tiers — working, episodic, semantic, procedural — from MemGPT paging to MemOS-style scheduling; priority queues for memory, with salience-, recency- and importance-scored admission control and score-based eviction; priority decay; provenance-chained updates; memory operations as callable tools; the silent-failure problem — a wrong eviction throws no exception, so memory-operation logs join the audit trail; LoCoMo, LongMemEval and BEAM. Retrieval: Corrective-RAG, structured retrieval over filings, stateful agents with Letta, chunking and citation discipline — every retrieved document treated as untrusted input.
LAB	Build the streaming covenant monitor over live 10-Q filings under an explicit context budget, with a priority-queue memory scheduler deciding what persists between filings — admission scores and eviction decisions logged as spans. Then plant an adversarial instruction inside a filing and confirm the agent reports it rather than obeying it.
DELIVERABLE	A covenant-monitor agent, its corpus configuration, a context-budget spec, a memory-scheduling policy (admission, priority, eviction), and a document-poisoning test result.
ASSIGNMENT	Graded lab 5 — the covenant monitor, its context budget, the memory-scheduling policy and the poisoning result.

6 LLM agents + RL layers in financial simulation

Where reinforcement learning genuinely improves an agent — and when the better upgrade is the harness, not the weights.

MON 16 NOV · 3 HOURS · PACHECO AZNAR

You ship: reproduced crew-versus-PPO results, one substitution of your own, an optional skill-distillation arm, and a limitations note

TOPICS	The supplied frozen crew-versus-PPO benchmark: fixed data, costs, risk limits and evaluation window; Sharpe, CVaR, compliance-violation rate, latency and cost; separating orchestration effects from policy effects, read from the shared telemetry rather than a spreadsheet. The 2026 agentic-RL landscape: ToolRL, ReTool and ARPO for tool-integrated training; AgentGym-RL for long-horizon multi-turn RL; verl as training infrastructure; reward hacking and reasoning collapse as the honest failure modes. Harness versus weights: skill distillation and memory as in-context self-improvement, against RL weight updates. Long rollouts checkpoint and resume on the durability substrate from session 2. The exercise is a controlled simulation, not evidence of production readiness.
LAB	Run and reproduce the supplied benchmark in FinRL-Meta; substitute one component of your own — your cost model, your risk limits, or a signal your desk actually uses — and rerun under the same frozen window. Optional third arm: a skill-distilled agent against the static agent, same window, same budget.
DELIVERABLE	Reproduced benchmark results, one controlled substitution, and a reviewer-facing limitations note stating which effects are orchestration, which are policy, and which the benchmark cannot see.
ASSIGNMENT	Graded lab 6 — five-seed reproduction, one isolated substitution, and a limitations note a reviewer can use.

7 Evaluation, governance & controlled pilots

Deciding whether the system has enough evidence for a controlled pilot — and enough operations to survive one.

MON 23 NOV · 3 HOURS · NOGUER I ALONSO

You ship: your evaluated capstone and a controlled-pilot evidence pack, with a live monitoring dashboard and an operations runbook

TOPICS	Harness validation separates replicate dispersion from instruction sensitivity; lock the design before viewing outcomes. Benchmarks include AgentBench, SWE-agent regression, and MCP-Bench and MCPMark for tool-use stress testing; risk metrics include CVaR, stress scenarios and failure rate. Controlled-pilot controls cover guardrails, kill-switches, budgets, monitoring, hand-back and EU AI Act documentation on the post-Omnibus timetable — Article 50 transparency duties applying from 2 August 2026, stand-alone high-risk obligations deferred to 2 December 2027 and embedded high-risk to 2 August 2028 — and why a deferred deadline is an argument for building the technical file now rather than later. GDPR access, rectification and erasure rights against long audit-trail retention for what an agent remembers. Change control for skills and memory, not only model versions. Production operations: online evaluation on sampled live traffic with human calibration of the judge, drift detection, agent SLOs and incident response.
LAB	Run the capstone through the evaluation harness and red-team it — prompts, documents, memory and skills — then build the monitoring dashboard on the cohort's own six weeks of telemetry, and rehearse the kill-and-resume drill end to end.
DELIVERABLE	An evaluated capstone plus the controlled-pilot evidence pack: evaluation results, red-team report, monitoring plan, live dashboard, operations runbook (SLOs, alerts, rollback, kill-switch drill) and a reviewer-facing limitations note.
ASSIGNMENT	Assessed as the final project — the capstone against its targets, with the evidence pack and the operations runbook.

CAPSTONE

What you take away

CAPSTONE & DELIVERABLES

Sessions 1 to 6 each end with a component. Session 7 assembles them into a finance-sector agent of your choosing, with the evidence to prove it meets its targets.

- ✓ Six lab systems and an evaluated capstone — seven working systems in all
- ✓ Cost, latency and failure instrumentation you can reuse — standards-based OpenTelemetry GenAI telemetry, not a bespoke logger
- ✓ Five written artefacts a reviewer can read on their own: durability posture note, memory-audit procedure, sandbox policy, memory-scheduling policy and operations runbook
- ✓ An audit and governance pack mapped to the EU AI Act as amended by the Digital Omnibus — the transparency duties already applying, and the high-risk documentation set built to its deferred timetable

ASSESSMENT & COMPLETION

- ✓ **Practical labs — 60%.** Six graded lab systems, one from each of sessions one to six — each marked out of 10 and due the Sunday after its session.
- ✓ **Final project — 30%.** The capstone agent against its safety and performance targets.
- ✓ **Participation — 10%.** Engagement and peer review of other crews' agents.
- ✓ The written artefacts are folded into the existing lab grades; the weighting is unchanged.
- ✓ Successful completion requires a final weighted mark of at least 70/100, attendance at six of the seven live sessions, and submission of all seven systems.
- ✓ Marks and certification are awarded individually.
- ✓ AIFI Professional Certificate upon successful completion.

Built for the review, not the demo

Most agent courses stop when the demo runs. This one goes on through what decides whether an agent survives a risk review.

+ Real market data

Labs run on SEC filings and market data. An agent that only works on a curated demo prompt fails visibly in the debrief.

+ Cost budgeted early

Every agent is instrumented for token cost and latency on standard spans. A crew that wins at ten times the price is reported as such.

+ Adversarial by default

Document poisoning, prompt injection, memory and skill poisoning, and red-teaming are built into the labs, not mentioned as a caveat.

+ Governance that maps

Kill-switches, resume drills, audit trails and an EU AI Act mapping on the post-Omnibus timetable that a second-line risk function can actually work with.

COURSE FACULTY

Miquel Noguer i Alonso, PhD, ARPM

Course director — leads sessions 1, 2, 5 and 7

Founder and Chief Science Officer of AIFI. Thirty years in quantitative finance, with prior roles at Andbank and UBS; teaches at ESADE Business School and at Cornell in New York City, with earlier positions at NYU and Columbia. Author of six books and more than 220 research papers.

David Pacheco Aznar

Faculty — Staq.io — leads sessions 3, 4 and 6

Co-author with the course director on AIFI's agentic-finance research, including the forthcoming Wiley volume and the agentic-AI modelling manifesto. He leads observability and decision tracing, tool-using and coding agents, and the RL benchmark — the same three sessions he taught in the 2025 edition of this certificate.

2,000 USD

Standard price per participant

1,600 USD Super Early Bird, to 28 Aug 2026

1,800 USD Early Bird, to 25 Sep 2026

Places are limited. Applications close 30 September 2026.

TUITION INCLUDES

- ✓ 21 core instructional hours across seven sessions
- ✓ Two faculty across the seven sessions
- ✓ Seven optional 60-minute office hours, Mondays after class
- ✓ Starter repository, notebooks, OTLP collector, checkpoint store, microVM images, evaluation harness and a dated reading list
- ✓ Six lab systems plus your evaluated capstone
- ✓ AIFI Professional Certificate upon successful completion

Bundle both October courses — with Advanced Claude and GPT for Finance: **2,300 USD** Super Early Bird, **2,500 USD** Early Bird, **2,600 USD** standard — savings of 100 USD, 200 USD and 400 USD respectively. Bundle applications close 30 September 2026. The two courses share the week of 12 October, which carries three sessions in five days on three separate days. **15% off** the tuition tier then in force for three or more participants from one firm. Group and bundle discounts are not cumulative.

OPERATING BOUNDARIES

Labs use public or synthetic data only. Participants must not upload client data, confidential documents or material non-public information, and must not connect the course environment to production systems. Where a session asks you to expose one of your own tools through MCP, use a redacted or reconstructed version of the function, not the production implementation or its data. The capstone is a controlled-pilot prototype with evidence designed to support review — not institutional approval, legal certification or permission to deploy against client assets. Regulatory dates are stated as they stand in July 2026 and are taught as a moving target, not as legal advice. Participants provide their own model-provider access and vector store; API and infrastructure charges are not included in tuition.

Registration & enquiries

info@aifinanceinstitute.com · www.aifinanceinstitute.com

