
A NEXTFI DIAGNOSTIC ENGAGEMENT · MAY 2026

Agentic AI Governance Review

Governance built for the agent, not the model.

PREPARED FOR

[Client Name Redacted]
Tier 2 US Commercial Bank

LEAD ADVISOR

Barry Eisenberg
Managing Principal

ENGAGEMENT PERIOD

April 2026
Three-Week Fixed Scope

ENGAGEMENT REFERENCE

AGR-2026-04-[redacted]
agentic.nextfiadvisors.com

REDACTED SAMPLE · PROSPECT REVIEW USE ONLY

Client name, deployment specifics, scores, and quantitative findings are illustrative. The structure, layout, depth of analysis, and tone are representative of the production deliverable.



What you're looking at.

A redacted, illustrative sample of the deliverable produced by a NextFi AGR engagement.

This is a redacted, illustrative sample of the deliverable produced by NextFi Advisors' Agentic AI Governance Review (AGR) — a fixed-scope, fixed-fee three-week engagement that scores an organization's agentic AI deployment across four governance lenses and produces a regulator-ready Findings Report and 90-Day Remediation Roadmap.

Client name, scores, quantitative findings, and deployment specifics have been replaced with redaction bars or generic placeholders. The structure, layout, depth of analysis, and tone are representative of the production deliverable. The Findings Report in the production engagement runs 40–60 pages; this sample includes ten content pages.

Contents

01	Executive Summary	03
02	Methodology — The Four-Lens Framework	04
03	Scorecard at a Glance	05
04	Detailed Findings — Control Plane	06
05	Detailed Findings — Regulatory & FS Risk	07
06	Control-Plane Map (Excerpt)	08
07	Regulatory Cross-Walk (Excerpt)	09
08	90-Day Remediation Roadmap	10
09	Next Steps	11

FORMAT
Fixed scope

TIMELINE
3 weeks

FEE
~\$10,000

OUTPUT
Regulator-ready



Reactive posture, with material gaps in agentic-specific failure modes.

Policies exist on paper; the control plane does not yet enforce them.

42

out of 100

REACTIVE

The Client's agentic AI deployment scored **42 / 100** on the NextFi Agentic AI Governance Review, placing it in the **Reactive** tier. Documented policies exist across most domains and the institution's traditional model risk framework is mature. However, governance is largely encoded in policy language rather than in the control plane itself, and the deployment is materially exposed to agentic-specific failure modes that classical model risk frameworks are not designed to address.

Headline findings

01 Permission enforcement lives in policy, not infrastructure.

Agentic deployments rely on prompt-level instructions and written policy to constrain tool access. No per-agent allow-list is enforced at the tool or MCP layer.

02 Memory and context governance are not yet operational.

Persistent memory is enabled by default without retention policy, redaction, or audit logging — creating a privacy and discovery exposure.

03 Shadow AI is the largest unmanaged regulatory surface.

Estimated **[redacted]**% of internal AI usage is on unsanctioned tools. No CASB or DLP-based discovery is in place; the institution is materially exposed to data exfiltration and AI-Act-relevant uses outside the governance perimeter.



Four lenses. One control plane.

A 24-dimension assessment designed so classical risk practitioners and agentic-native engineers both see their domain represented.

The AGR scores the deployment across four governance lenses, each with six diagnostic dimensions. Together they produce a 24-dimension assessment scored 0–4 per dimension. Lens scores roll up to /25; the total to /100.

PROPRIETARY LENS

Control Plane Architecture

01

Governance encoded in infrastructure, not policy: permission engines, context pipelines, tool boundaries, memory, delegation, observability.

DEFENSIBILITY

Regulatory & Standards Alignment

02

NIST AI RMF, EU AI Act, SR 11-7, EU DORA, SEC cybersecurity disclosure, third-party AI risk, acceptable-use enforcement.

DOMAIN COVERAGE

Financial-Services Risk Domains

03

Autonomous action, consumer protection and fair lending, BSA/AML and sanctions, AI-enabled fraud and deepfakes, data sovereignty, concentration risk.

BEYOND CLASSICAL AI RISK

Agentic-Specific Failure Modes

04

Prompt injection, tool-use abuse, runaway loops, memory poisoning, multi-agent cascading failure, shadow AI and unsanctioned agents.

Scoring. Each dimension is scored on a 0–4 scale anchored to control maturity, where 0 indicates no awareness or control and 4 indicates the control is enforced at the infrastructure layer, monitored continuously, and audit-ready.



Lens Scorecard

Four governance lenses, normalized 0–25, weighted to a 0–100 total.

42

out of 100

Reactive

TIER 2 OF 4

Policies exist on paper; infrastructure does not. Governance is documented but not enforced at the control plane. The shadow-AI gap is likely significant.

Control Plane Architecture

Proprietary lens · 30% of total



7 / 25

Regulatory & Standards Alignment

Defensibility · 25% of total



14 / 25

Financial-Services Risk Domains

Domain coverage · 25% of total



16 / 25

Agentic-Specific Failure Modes

Beyond classical AI risk · 20% of total



5 / 25

0 6 13 19 25

Reading this page. Each lens is scored 0–25 from the diagnostic responses, then weighted by its share of the total (Control Plane 30%, Regulatory 25%, FS Risk 25%, Agentic Failure 20%). The institution's weakest lenses are **Agentic-Specific Failure Modes** (5 / 25) and **Control Plane Architecture** (7 / 25) — driven by the absence of an enforced permission engine, missing memory-write governance, and no automated shadow-AI discovery. Detailed findings begin on page 6.



Permission, context, and memory live in policy, not infrastructure.

Detailed findings — Control Plane Architecture lens.

FINDING 1.1 · PERMISSION ENGINE

HIGH SEVERITY

Tool access governed by prompt instructions, not enforced at the call layer.

Current 1 / 4 **Target** 4 / 4 **Effort** ~6 weeks **Owner** CISO / AI Lead

OBSERVATION

Across the three deployments reviewed, the agent's ability to call external tools is governed by the system prompt and a written acceptable-use policy. There is no enforcement layer between the agent and the tool: any tool the agent can reference, it can invoke. Credentials scoped to the deployment grant broad read/write access to the underlying systems.

WHY IT MATTERS

Prompt-level instructions are not enforcement; they are suggestions. A prompt-injected document, a poisoned retrieval result, or a deliberately crafted user request can override the system prompt without leaving an enforcement-layer signal. The institution's existing model-risk framework does not contemplate this surface because it inherits assumptions from traditional model deployments.

RECOMMENDATION

Adopt an MCP-layer permission engine that enforces per-agent, per-tool, per-data-classification access policy independent of the agent prompt. Anchor the design to NextFi's published guidance in "MCP: The USB-C for AI Integrations" and the Agentic Architecture, Annotated brief. Implementation sequence is detailed in Roadmap item R-01.

FINDING 1.3 · MEMORY GOVERNANCE

CRITICAL

Persistent memory enabled without retention, redaction, or audit policy.

Current [redacted] / 4 **Target** 4 / 4 **Effort** [redacted] weeks **Owner** Data Governance

OBSERVATION

Full finding text available in the unredacted Findings Report. Production deliverable includes evidence references, control-plane layer mapping, and remediation owner.



Defensibility gaps cluster around the EU perimeter and shadow AI.

Detailed findings — Regulatory & Standards Alignment and FS-Risk Domains lenses.

FINDING 2.2 · EU AI ACT

HIGH SEVERITY

EU AI Act risk-tier classification not performed for any in-scope deployment.

OBSERVATION

Although the Client serves EU clients through [redacted], no formal classification of the agentic deployment under EU AI Act risk tiers has been performed. The deployment may include uses that fall under the high-risk category, triggering conformity assessment, post-market monitoring, and transparency obligations.

RECOMMENDATION

Conduct a formal EU AI Act classification within 30 days. Where high-risk uses are identified, [redacted] a documented conformity assessment plan within 60 days. Detail in Roadmap item R-04.

FINDING 2.6 · SHADOW AI

CRITICAL

Shadow AI exposure is the largest unmanaged regulatory surface.

OBSERVATION

Discovery scans conducted during the engagement indicate that approximately [redacted]% of internal AI usage occurs on tools that are not part of the institution's sanctioned AI estate. Categories include [redacted]. There is no CASB or DLP-based discovery in place. Acceptable-use policy is not technically enforced.

RECOMMENDATION

Stand up a continuous AI-discovery capability (CASB or AI-specific tooling) within 45 days. Establish a sanctioned-by-default workflow with documented exception path. Detail in Roadmap item R-05.

FINDING 3.4 · AI-ENABLED FRAUD

HIGH SEVERITY

AI-enabled fraud and deepfake defense is awareness-only.

OBSERVATION

Existing fraud-prevention controls are designed for [redacted] and have not been updated to address voice cloning, synthetic-identity injection, or agentic social engineering. Liability is shifting toward the institution that processes the fraudulent action.

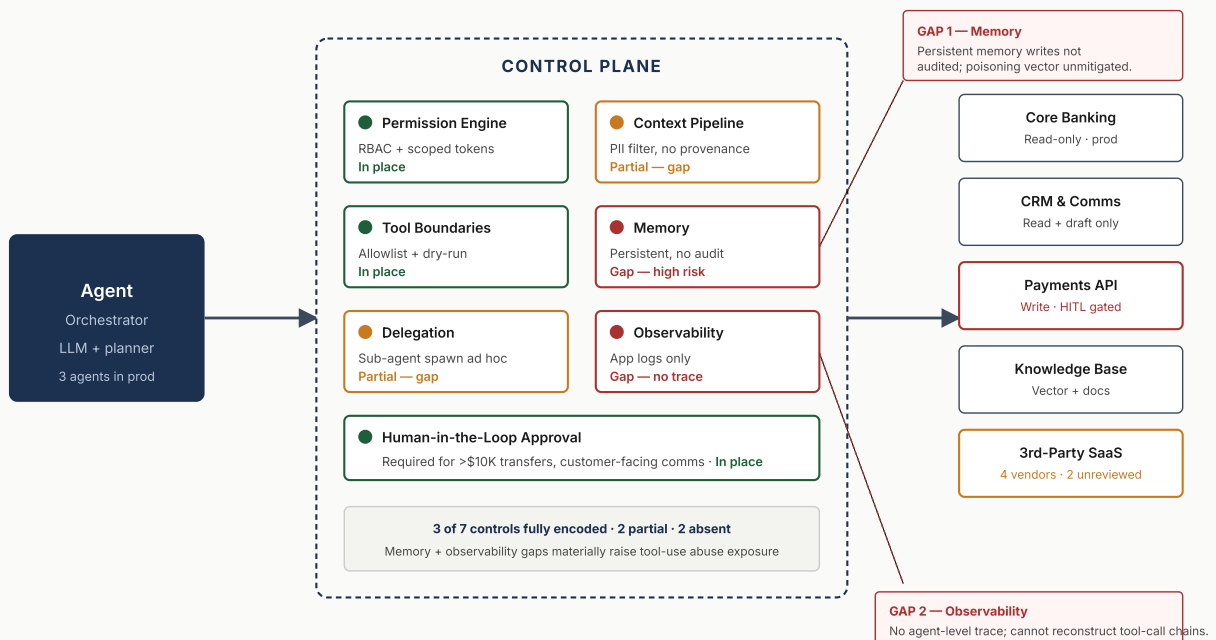
RECOMMENDATION

Deploy liveness checks and behavioral biometrics on customer-facing voice and document channels within 90 days. Detail in Roadmap item R-07.



Control-Plane Map

Where governance is — and is not — encoded in the institution's agent infrastructure.



● Control encoded & tested ● Partial — documented but not enforced ● Absent — material gap

Annotations flag findings detailed on pages 6–7.

Regulatory cross-walk — NIST AI RMF excerpt.

Each finding is mapped to specific subcategories of NIST AI RMF, EU AI Act articles, SR 11-7 sections, and DORA chapters.

CATEGORY	SUBCATEGORY	DESCRIPTION	STATUS	FINDING
GOVERN	1.1	Legal and regulatory requirements understood	PARTIAL	2.1, 2.2
GOVERN	1.4	Risk management processes integrated with broader enterprise risk management	DOCUMENTED	—
GOVERN	2.1	Roles and responsibilities for AI risk are documented and communicated	PARTIAL	1.6
GOVERN	4.1	Policies for AI use are established and enforced	GAP	2.6, 1.1
MAP	1.1	Context of AI system use is documented	DOCUMENTED	—
MAP	2.3	Scientific integrity and TEVV considerations addressed	PARTIAL	1.5
MAP	5.1	Likelihood and magnitude of impact identified	PARTIAL	2.2
MEASURE	2.7	AI system security and resilience are evaluated	GAP	4.1, 4.2
MEASURE	2.10	Privacy of AI system is examined and documented	GAP	1.3
MANAGE	1.3	Risk responses are developed and prioritized	PARTIAL	1.1, 1.3

Excerpt only. The production deliverable includes the full cross-walk (47 rows across NIST AI RMF, EU AI Act, SR 11-7, and DORA), each row evidence-linked to specific findings.

90-Day Implementation Roadmap

Prioritized to close the two material gaps (Memory, Observability) before scaling agent count.

WORKSTREAM	W1	W2	W3	W4	W5	W6	W7	W8	W9	W10	W11	W12
PHASE	FOUNDATION				BUILD				VALIDATE			
Control Plane Memory audit · agent tracing			M1									
		Memory write log (url)	Agent trace instrumentation						Tabletop replay test			
Regulatory SR 11-7 + EU AI Act mapping			Control crosswalk		Evidence binder draft					Internal audit r		
FS Risk Integration BSA/AML · consumer protection								M2				
			Taxonomy exte	Agent-action risk-tagging in i	Monitor							
Agentic Failure Prompt-injection · shadow-AI												M3
		Shadow-AI dis	Prompt-injection test suite			Tool-use anomaly rule			Red-team exercise			
Governance & Reporting Board cadence · disclosure												
		Charter + RACI				Monthly metrics pack					Board readout	

M1 · End of week 4

Memory writes audited; trace instrumentation deployed across all 3 production agents.

M2 · End of week 8

Shadow-AI sweep complete; prompt-injection test suite passing on baseline policy.

M3 · End of week 12

Red-team exercise complete; first board readout delivered; SR 11-7 evidence binder reviewed.



From findings to a governed deployment.

Three concrete next steps that move the deployment from Reactive toward Managed within one quarter.

The Agentic AI Governance Review concludes with this Findings Report. The institution has three immediate options for moving from Reactive toward a defensible, governed posture.

01 Adopt the 90-Day Roadmap as-is.

The roadmap is sized for the institution's existing AI, security, and risk teams. It does not require external implementation support and is designed to close the two highest-severity findings (Memory Governance, Shadow AI) within the first 30 days.

02 Engage NextFi for selected implementation support.

Per-workstream support is available for the Permission Engine implementation (R-01), Memory Governance design (R-02), the Shadow AI discovery program (R-05), and the EU AI Act conformity assessment (R-04). Each is scoped as a fixed-fee, time-boxed engagement.

03 Schedule a 90-day re-score.

NextFi recommends a re-score at the end of the 90-day window to validate uplift, refresh the regulatory cross-walk, and update the control-plane map for any infrastructure changes. Re-scores are delivered against the same 24-dimension framework so progress is directly comparable.

Engagement specifications

FORMAT	TIMELINE	INVESTMENT
Fixed scope, fixed fee	Three weeks	~\$10,000*

**Typical engagements range between \$5k and \$20k depending on complexity*

FROM FINDINGS TO A GOVERNED DEPLOYMENT

A report like this, for your stack.

The 24-question NextFi diagnostic returns a free scorecard in about ten minutes. From there, an Agentic AI Governance Review produces the full findings report — mapped to your control plane, regulatory exposure, and a 90-day remediation roadmap.

agentic.nextfiadvisors.com →