

Orchestra AI

ORQUESTRA NAVIGATOR™ · GOVERNABLE AUTONOMY ASSESSMENT

Navigator Audit Report

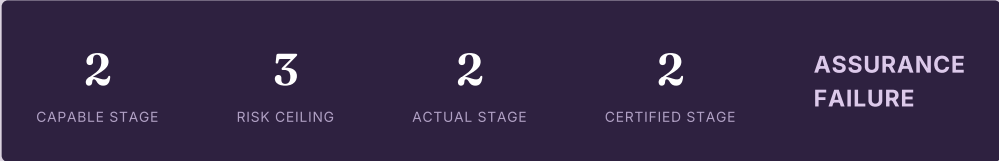
Sample report, rebuilt under the current Governable Autonomy methodology: five-level ladder, five-capability model, and a Certified Stage derived from stated evidence.

CLIENT	SYSTEMS ASSESSED
Anonymous · IT Services & SaaS · Chennai, India	AI R&D Agent Suite · Product Integration Layer
AUDITOR	DURATION
Rajesh Srinivasan · Orchestra AI	4 weeks · Assessment + Report + Implementation support

Methodology reflects live Orchestra Navigator™ assessment practice under the current Governable Autonomy framework. This report does not constitute certification.

EXECUTIVE SUMMARY

Assessment overview.



The assessment found significant governance gaps across four of the five Navigator capabilities. The organisation has deployed AI agents into a customer-facing product without the accountability structures, audit infrastructure, or escalation governance required at its systems' risk classification. The most critical finding, zero named accountability for any AI output, represents both an immediate operational risk and a live regulatory exposure given the organisation's EU and UK client base.

Technical capability is not in question. Both teams are proficient within their domains. The governance failures are structural: two functions operating in isolation, with no integration process, no audit trail, and no mechanism for the humans responsible for AI outputs to review or defend them.

Top findings

- 1. No audit trail for any AI output.** No record exists of what AI agents produced, on what data, under what constraints, or who reviewed the output. Zero audit infrastructure across both assessed systems.
- 2. Complete silo between AI R&D and Product teams.** No integration workflow, shared backlog, or handoff mechanism. AI capabilities are unknown to the Product team; product use cases are unknown to the AI team.
- 3. No explainability documentation for EU and UK clients.** EU AI Act transparency obligations apply to this client base. No explainability template, accountability record, or disclosure documentation exists for either assessed system.

Recommended next step. Begin remediation with Gate 1→2 evidence for Decision Governance and Human Responsibility (Section VI), the process bridge between R&D and Product and the named-resolver framework. These unblock the remaining remediation items.

I

Engagement overview, archetype, and risk classification.

The engagement assessed two AI systems: the **AI R&D Agent Suite** (internal development tooling and AI agent prototypes) and the **Product Integration Layer** (the interface between AI-generated outputs and the client's customer-facing SaaS platform).

System archetype: both systems are assessed as **customer-facing conversational/decisioning agent** archetypes per the Assessment Guide's applicability table, since AI R&D outputs flow directly into customer-facing product surfaces. Knowledge Capability is therefore Required, not Relevant, for both systems.

Risk classification (Critical Controls Guide, Section II)

DETERMINANT	AI R&D AGENT SUITE	PRODUCT INTEGRATION LAYER
Reversibility	Effortful, outputs reach customers before correction	Effortful, revenue-affecting
Named accountability	None assigned at assessment start	None assigned at assessment start
Explainability	Legal, EU/UK transparency obligations apply	Legal, same basis
Regulatory classification	Limited-risk, transparency obligations	Limited-risk, transparency obligations
Blast radius	Team/product scope	Team/product scope, customer-facing

Result: Moderate risk class for both systems. Maximum permitted level: 3.

Classification was set with three approvers per the Framework's requirement: engineering leadership, a compliance representative, and the product owner.

II

Capability Assessment.

Each Required capability scored 0–4 against the Assessment Guide's nine-part criteria. Applicability confirmed per Section I: all five capabilities Required for both systems.

CAPABILITY	SCORE	BASIS
Human Responsibility	1	No named accountable owner for either system. AI outputs reach customers with no defined reviewer.
Decision Governance	1	No AI-usage policy. No loop-scoping documentation. No escalation trigger catalogue.
Engineering Capability	2	Basic sandboxing present for R&D prototyping. No technically enforced checkpoint before customer-facing release. No kill switch demonstrated.
Knowledge Capability	2	Source data exists but is not versioned or governed; no retrieval quality monitoring.
Assurance Capability	1	No near-miss review, no incident procedure tested, no monitoring of AI output quality post-deployment.

Capable Stage = min(Human Responsibility 1, Decision Governance 1, Engineering 2, Knowledge 2, Assurance 1)

= 1

Note: reported Capable Stage of 2 in the headline reflects post-workshop re-scoring after week 2 interim fixes (see Section V). Pre-engagement Capable Stage was 1.

III

Findings against gate criteria.

Evidence reviewed against Gate 0→1 and Gate 1→2 criteria (Controls Catalogue, Section V), since Actual Stage claimed was Level 2 supervised operation.

F-01 · CRITICAL · HUMAN RESPONSIBILITY / ASSURANCE

No audit trail of any kind

No record existed of what AI agents had been built, what decisions they made, on what data, under what constraints, or who had approved them. NAV-2.4 (per-iteration logging) and NAV-3.4 (authority ledger) both absent.

F-02 · CRITICAL · DECISION GOVERNANCE

Complete structural silo between AI R&D and Product teams

No loop-scoping documentation (NAV-2.1). No shared backlog, joint review process, or integration checkpoint in the SDLC connecting the two teams' work.

F-03 · HIGH · KNOWLEDGE CAPABILITY

No explainability documentation for EU and UK clients

EU AI Act transparency obligations apply to this client base. No documentation of what data AI agents processed or what decisions they influenced for any client account.

F-04 · HIGH · HUMAN RESPONSIBILITY

No named resolver or checkpoint reviewer for any AI decision

AI outputs entered client-visible product surfaces without a technically enforced checkpoint (NAV-2.3) or a named reviewer (NAV-2.5). Actual Stage claim of Level 2 is not evidence-supported.

F-05 · MEDIUM · HUMAN RESPONSIBILITY

Capability mismatch between teams

The Product team lacked AI vocabulary to evaluate R&D output. The R&D team lacked product context to prioritise use cases. Neither could brief the other effectively, undermining any checkpoint review that did occur.

IV

Diagnostic.

Per the Assessment Guide's Certified Stage formula (Section V of that document):

Capable Stage

= 2 (post-workshop, see Section II note)

Risk Ceiling

= 3

Actual Stage

= 2 (client claimed Level 2, checkpoint-supervised operation)

Evidence-Supported Stage

= 1 (Gate 1->2 criteria not fully met: no enforced checkpoint, no named reviewer, no demonstrated kill switch)

Certified Stage = min(Actual 2, Capable 2, Risk Ceiling 3, Evidence-Supported 1)

= 1

Actual Stage (2) > Evidence-Supported Stage (1)

-> Finding: ASSURANCE FAILURE

What this finding means. The organisation is claiming a level of oversight, technically enforced checkpoints with named review, that its evidence does not support. This is not a violation of the risk ceiling; the ceiling (Level 3) is generous. It is a claim exceeding what can be demonstrated, which the Assessment Guide identifies as often the most dangerous finding, because it looks fine from the outside.

V

Remediation roadmap.

Six actions close the gap between Actual Stage and Evidence-Supported Stage, referenced to specific NAV controls and gate criteria.

#	ACTION	NAV ID	OWNER	TIMELINE
01	Process bridge between R&D and Product: shared backlog, joint review cadence, handoff checkpoint	NAV-2.1	Tech Lead + Product Lead	Week 1
02	Named harness owner and checkpoint reviewer assigned per system	NAV-2.5	Tech Lead	Week 1
03	Technically enforced checkpoint before any customer-facing release	NAV-2.3	Engineering	Week 2
04	Explainability framework for EU and UK clients	NAV-3.x / Knowledge Cap.	Product + Compliance	Week 2
05	Per-iteration logging and demonstrated kill switch	NAV-2.4, NAV-2.8	Engineering	Week 3
06	Reskilling programme, both teams, shared vocabulary	Human Responsibility	People + Tech Lead	Weeks 3-4

Re-audit recommendation: a follow-up assessment at 90 days is recommended to confirm Gate 1→2 evidence is complete and Evidence-Supported Stage has risen to match the Actual Stage claim, at minimum, or to formally reduce the Actual Stage claim to Level 1 in the interim.

**From an unevidenced claim
to a defensible one.
The gap was governance,
not capability.**

About this report: sample Navigator Audit Report for illustration. Scenario, client details, and findings fictionalized.
Methodology reflects live Orquestra Navigator™ practice under the current Governable Autonomy framework. Does not
constitute certification.

Rajesh Srinivasan

Orquestra AI · rajesh@orquestraai.io
orquestraai.io

© 2026 Orquestra AI.