

Execution Governance for Agentic AI in Financial Infrastructure

Authorization before execution. Enforcement during execution.
Cryptographic proof after execution.

11 AI AND BLOCKCHAIN DEVELOPMENT LLC

11/11 AI Research Division

*Technical white paper accompanying the response to the FSB consultation report
Sound Practices for Responsible Adoption of Artificial Intelligence (AI), 10 June 2026*

30 N Gould St Ste R, Sheridan, WY 82801 • quantum@11aiblockchain.com • 11aiblockchain.com

July 2026

Execution Governance™ • Governed Execution™ • EA-11™ • Execution Arithmetic™ • Execution Evidence State™ • Patent Pending

Contents

1. Executive Summary
2. Introduction
3. Problem Statement
4. The Current AI Governance Landscape and Its Boundary
5. Why Agentic AI Changes Financial Risk
6. The Execution Governance Architecture
7. Financial Infrastructure Applications
8. Governance APIs, Implementation Model, and Open Standards
9. Mapping to the FSB Sound Practices
10. Recommendations
11. Conclusion
12. References

Basis of this paper

Every architectural capability described in this paper traces to published, DOI-registered doctrine and specifications of the 11/11 AI Research Division, to open-source reference tooling, or to publicly accessible verification endpoints. The complete catalog, with DOIs, repositories, and endpoints, is provided in the companion document, 11/11 AI Public Technical Evidence Appendix. This paper asserts no market adoption, no regulatory endorsement, and no third-party certification.

1. Executive Summary

Financial institutions are moving from AI systems that predict and recommend to AI systems that act. Agentic AI, defined by the Financial Stability Board as autonomous systems “capable of planning, reasoning, and executing complex high-level goals (comprising multiple tasks) independently,” introduces a risk surface that established AI governance does not reach: the individual execution. A model can be validated, documented, and overseen according to every current sound practice and still, at runtime, attempt an action that is outside its mandate: a transfer above threshold, a query beyond scope, a message to a system it should not touch. In financial infrastructure, where settlement finality can make actions irreversible, this gap is material to financial stability.

This paper describes execution governance: a runtime governance layer that authorizes, enforces, and proves AI execution. The architecture rests on a triad. First, authorization before execution: every proposed action is bound to a subject, an action, a resource, and a context, and is evaluated against the institution's active policy set before it is permitted to run. Second, enforcement during execution: an action proceeds only in the presence of a valid authorization artifact, and the default in the absence of one is denial (fail-closed). Third, cryptographic proof after execution: every decision, permitted or denied, produces a signed, hash-chained, independently verifiable record.

The paper positions execution governance as infrastructure, not as a model or a model-safety technique. It complements model governance, data governance, and risk governance; it does not replace them. All architectural claims trace to a published, DOI-registered doctrine and specification corpus (the RFC-EG series and related doctrine of the 11/11 AI Research Division), to open-source reference tooling, and to a live control plane whose signed decisions are publicly verifiable without credentials. The paper maps the architecture to each of the FSB's twelve proposed sound practices for responsible AI adoption and closes with recommendations for standard-setters and supervisors. It makes no claims of market adoption, regulatory endorsement, or third-party certification; it argues from architecture and public evidence.

2. Introduction

Three families of governance dominate current practice in financial AI. Model governance addresses how models are developed, validated, and monitored, typically under model risk management regimes descended from supervisory guidance on model risk. Data governance addresses the quality, provenance, lineage, and permitted use of the data on which models train and operate. Risk governance addresses how AI risk is identified, owned, measured, and escalated across the three lines of defence. The FSB's June 2026 consultation report organizes these strands into twelve sound practices covering organization-wide governance, the AI lifecycle, and cyber and third-party risk.

These families share a common object: they govern the system and its lifecycle. They ask whether the model was built well, whether its data is sound, whether its risks are owned. Agentic AI adds a fourth question that lifecycle governance cannot answer from inside its frame: is this specific action, about to be taken by this specific agent, on this specific resource, at this specific moment, authorized? Answering that question at runtime, uniformly and provably, is the province of execution governance.

The term is used here in a precise sense. Execution governance is a runtime control layer, external to the AI system it governs, that (i) evaluates every proposed action against explicit policy before the action runs, (ii) enforces the resulting decision such that ungoverned execution is not possible within the governed perimeter, and (iii) emits cryptographic evidence of every decision. The layer is indifferent to the internals of the AI that proposes the action. It governs conduct, not cognition.

Scope and sources. This paper is not about AI ethics, alignment, or the safety of large language models. It concerns runtime governance, execution assurance, and operational trust in financial infrastructure. Every capability described is published in the open corpus catalogued in the accompanying Technical Evidence Appendix: DOI-registered doctrine and specifications, open-source verification tooling, and live public proof endpoints. Where the paper describes the operating characteristics of 11/11 AI's control plane, those characteristics are observable on its public surfaces.

3. Problem Statement

In a conventional deployment, an agent that holds credentials to an API, a database, or a payment interface can exercise the full authority of those credentials the moment it decides to act. Whatever governance preceded deployment, at the instant of execution nothing independent stands between the agent's decision and its effect. The FSB consultation report describes the consequences with unusual directness: agents can take “illegal, unethical, or unauthorised actions without human approval or oversight,” and “[o]verriding, redressing, or remediating these actions can be difficult or impossible for humans.”

The report equally identifies why the traditional compensating control does not scale: “AI agents pose a distinct challenge for human oversight, given the impracticality of real-time human monitoring of agent decisions as their use scales,” and an agent “can take hundreds of intermediate steps in pursuit of its goals and make errors in any of those steps.” Post hoc review presumes reversibility; in payments and settlement, finality removes it. Logging presumes trustworthy records; logs written by the system under examination are neither signed nor tamper-evident. Performance thresholds presume measurable outputs; for agents, the report notes, what must be assessed is “not only whether the AI agent achieved its overall objective but also whether the actions it took to do so were appropriate.”

The problem, stated as an engineering requirement: financial institutions deploying agentic AI require a control point at the execution boundary that is (a) independent of the governed system, (b) evaluated before the action takes effect, (c) fail-closed under uncertainty or failure, (d) uniform at any scale of agent activity, and (e) evidentiary, producing records that third parties can verify. No combination of model, data, and risk governance supplies this control point, because all three operate on the system and its lifecycle rather than on the individual execution.

4. The Current AI Governance Landscape and Its Boundary

The FSB's twelve proposed sound practices are the current reference point for the landscape, and this paper takes them as given. They group as follows: organization-wide governance (Sound Practices 1 to 4: strategic direction and oversight; governance and accountability; incorporation of AI risks into the risk management framework; organisational adaptability), AI lifecycle management (Sound Practices 5 to 10: materiality and risk assessment; selection; data governance; explainability and transparency; performance management; human oversight), and cross-cutting risk (Sound Practices 11 and 12: cyber and ICT risk management; third-party AI risk management), applied across a seven-stage lifecycle from inception to retirement and calibrated by proportionality. Related instruments occupy the same space: model risk management regimes, the NIST AI Risk Management Framework, operational resilience and third-party risk guidance, and the FSB's 2017 and 2024 assessments of AI's financial stability implications. What all of these share is their object of control: the system and its lifecycle.

The following comparison is definitional, not critical. Lifecycle governance is necessary; the point is only that its object of control is different, and that agentic AI makes the difference operationally significant.

Dimension	Traditional AI governance	Execution governance
Object of control	The model or system and its lifecycle (inception through retirement).	The individual execution: one action, by one subject, on one resource, in one context.
When control applies	Before deployment and periodically thereafter (validation, review, monitoring).	At the moment of each action: before, during, and after execution.
Mechanism	Documentation, validation, testing, human oversight, committee review.	Policy evaluation, authorization artifacts, runtime admission, signed evidence.
Failure default	Typically fail-open: absent intervention, the deployed system continues to act.	Fail-closed: absent valid authorization, the action does not run.
Evidence	Reports, inventories, logs maintained by or near the governed system.	Cryptographic receipts: signed, hash-chained, verifiable by external parties.
Scaling behavior	Human-anchored controls degrade as decision volume grows.	Deterministic evaluation is uniform at any volume of agent activity.
Primary question	Was the system built, validated, and overseen properly?	Was this action authorized, and can that be proven?

Table 4-1. Traditional AI governance and execution governance compared. The layers are complementary; each governs what the other cannot.

A second comparison isolates the model-governance case, since model risk management is the regime financial institutions most naturally extend to AI.

Aspect	Model governance	Runtime governance
Grounding regime	Model risk management guidance; validation and challenger practice.	Authorization, admission control, and evidence specifications (for example, the RFC-EG series).
Unit of review	A model version and its documentation.	An execution and its authorization decision.
Cadence	Periodic: at development, validation, and re-validation.	Continuous: every action, every time.
Detects	Conceptual unsoundness, drift, degraded performance.	Out-of-policy actions, scope creep, unauthorized access, anomalous execution patterns.
Cannot detect	A validated model acting outside its mandate at runtime.	Latent statistical flaws in the model itself.
Analogy	Licensing a driver.	Traffic law, enforced on every journey, with a signed record of each.

Table 4-2. Model governance and runtime governance govern different failure modes; agentic AI requires both.

5. Why Agentic AI Changes Financial Risk

5.1 Execution risk

Execution risk is the risk that an AI-initiated action takes effect without independent authorization at the moment of execution. It is distinct from model risk (the model may be flawless), from data risk (the data may be sound), and from operational risk in the classical sense (no process failed; the process simply never contemplated an autonomous actor). Execution risk is a function of three variables the FSB report itself identifies as risk factors: the agent's autonomy, its access to systems and data, and its scope of action. Where all three are high and no execution-boundary control exists, a single erroneous decision propagates at machine speed into systems of record.

5.2 Autonomous decision risk

Agents plan. A goal decomposes into steps the institution never individually reviewed, and the report notes that agents “can also dynamically set or modify their objectives based on what they learn from interacting with their external environments.” Reward hacking, goal misalignment, and prompt-injection or memory-poisoning attacks all share one signature: the agent takes an action that no policy would have approved. Controls that examine outputs after the fact see the consequence; a control that evaluates each action before it runs sees the cause, and can stop it.

5.3 Multi-agent risk

The report warns that agents “may collude with other AI agents” and that “multiple AI agents may interact in unanticipated ways that create unforeseen behaviours or risks.” Multi-agent systems compound execution risk in two directions: delegation chains obscure which principal authorized what, and emergent interaction produces action sequences no single agent's oversight regime anticipated. Governance at the execution boundary is the natural control here because it is agent-agnostic: every action by every agent, including sub-agents and delegated tasks, passes the same authorization gate, and lineage records reconstruct the delegation chain after the fact. Published doctrine addresses inter-agent authorization, lineage, and audit directly (DOC-EG-003).

5.4 Financial infrastructure risk

Financial infrastructure concentrates every aggravating factor: irreversibility (settlement finality), speed (machine-tempo execution), interconnection (an erroneous instruction propagates across participants), and systemic reach (payment systems, CCPs, and custodians are single points whose disruption is a financial stability event, as the FSB's 2024 assessment of AI-related vulnerabilities recognizes). The FSB report already recommends “restricting direct agent access to payment systems (requiring human intermediation for execution)” and dual authorization above value thresholds. Execution governance is the generalization of those recommendations into infrastructure: the checkpoint is architectural, uniform, and provable rather than procedural.

5.5 The financial stability perspective

The FSB's 2024 assessment identified four AI-related vulnerabilities with financial stability potential: third-party dependencies and service provider concentration; market correlations; cyber risks; and model risk, data quality, and governance. Agentic execution sharpens each. Concentration means many institutions may run agents built on the same foundation models; correlated model behavior becomes correlated action, at machine tempo, unless institution-level policy gates diversify what those agents are permitted to execute. Cyber compromise of an agent matters in proportion to what the agent can do; an external, fail-closed authorization layer converts a compromised agent from an insider with credentials into a supplicant whose requests are policy-checked. And where the consultation report warns of homogeneous behaviors amplifying herding and procyclicality, execution-layer telemetry (authorization and denial flows) is the natural observable through which institutions,

and potentially macroprudential authorities, could detect correlated agent activity while it is still deniable rather than after it is settled. These systemic applications are directions rather than published capabilities, but they follow from the architecture's placement at the boundary where AI activity becomes financial fact.

Risk (FSB Section 1.3)	Mechanism	Execution governance control
Unauthorized actions	Agent acts beyond mandate; override after the fact difficult or impossible.	Pre-execution authorization; actions outside policy are denied before effect (fail-closed).
Erroneous actions	Goal misalignment, reward hacking, errors across hundreds of intermediate steps.	Per-step admission control; every intermediate action is individually authorized and recorded.
Data breaches	Agent exposes or manipulates sensitive data, maliciously induced or accidental.	Resource-scoped authorization artifacts; least privilege enforced per execution, not per session.
Disruption of connected systems	Compromised or malfunctioning agent reaches APIs, databases, other agents.	External enforcement point; a manipulated agent cannot self-authorize beyond policy.
Agent collusion / emergent interaction	Multi-agent behavior no single oversight regime anticipated.	Uniform gate across all agents and sub-agents; lineage reconstructs cross-agent causality.
Inadequate human oversight	Real-time human monitoring impractical at scale.	Humans govern the policy set and exceptions; the control plane applies policy to every execution.
Memory poisoning / novel evasion	Injected data shifts agent behavior over time.	Behavioral drift surfaces as denied executions; evidence chain supports root-cause analysis.

Table 5-1. Agentic AI risk matrix: the FSB's identified agentic risks mapped to execution-boundary controls.

6. The Execution Governance Architecture

This section describes the architecture as published in the open corpus: the foundational doctrine (DOI 10.5281/zenodo.20252295 and 10.5281/zenodo.20252639), the category definition (EG-001), and the RFC-EG specification series. The reference architecture is specified in RFC-EG-1000 and the control-plane decomposition in RFC-EG-0020.

Figure 6-1. The Execution Governance Triad

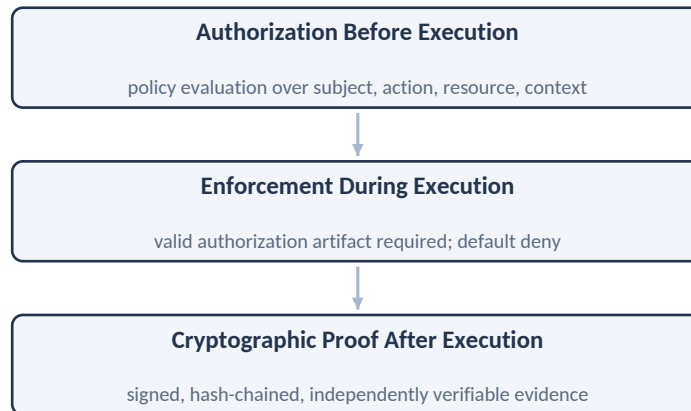


Figure 6-1. The triad. Every element of the architecture serves one of these three phases.

6.1 Authorization before execution

Every proposed action is expressed as a four-part binding: a subject (the identity of the agent or system proposing the action), an action (the operation), a resource (the target), and a context (the environment and conditions). The binding is evaluated against the institution's active policy set before the action is permitted to run. The output is a governance decision; where the decision permits execution, it is embodied in an authorization artifact, a signed data object specified in RFC-EG-0100 that downstream systems require before executing. Subject identity is carried by an identity token profile specified in RFC-EG-0120, addressing what the FSB report calls dynamic identity and access management for AI agents; where hardware provenance matters, a device attestation chain is specified in RFC-EG-0220. Evaluation is deterministic: the same request against the same policy set yields the same decision, which makes authorization behavior testable in advance and reproducible in audit, in contrast to oversight mechanisms that depend on human availability or model judgment.

Figure 6-2. Authorization Pipeline

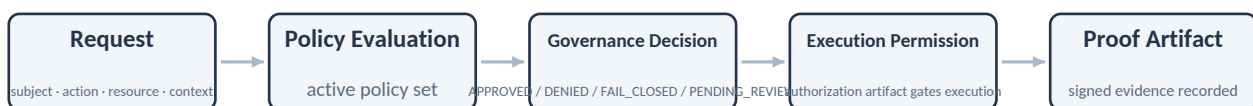


Figure 6-2. The authorization pipeline as documented for the public control plane. Decision states are APPROVED, DENIED, FAIL_CLOSED, and PENDING_REVIEW.

6.2 Enforcement during execution

Authorization without enforcement is advice. The enforcement property, specified as runtime admission in RFC-EG-0200, is that an action proceeds only when a valid authorization artifact is present; without it, execution is stopped. Enforcement is fail-closed: if authorization is missing, invalid, expired, or cannot be evaluated (including under control-plane degradation), the default outcome is denial. Fail-closed is the property that

converts written risk appetite into physical constraint. A prohibited use case under Sound Practice 1 ceases to be a policy statement that an agent might violate and becomes a class of actions that cannot execute within the governed perimeter. It is also the property that answers the report's kill-switch consideration architecturally: rather than an added mechanism that humans must reach in time, constraint is the resting state, and execution is the exception that policy explicitly grants. Graduated degradation, from autonomous operation to human approval checkpoints to full stop, is expressed as policy rather than engineered per system.

6.3 Proof after execution

Every decision, permitted and denied alike, produces evidence. A denial is not a silent failure; it is a signed, persisted record carrying a policy-violation reason. The published cryptographic profile names the primitives at each layer: Ed25519 signatures over authorization artifacts, SHA3-512 for audit-chain hashing, and BLAKE2b-512 for secondary evidence chaining. For long-lived evidence, a post-quantum signature profile (RFC-EG-0301) specifies ML-DSA-65 and SLH-DSA alongside Ed25519; the public evidence endpoint operates this profile today, and the rationale for post-quantum protection of long-retention AI evidence is treated in doctrine (DOC-EG-005, on the cryptographic half-life of AI artifacts). Financial records routinely outlive a decade; evidence that must still verify then should be signed under assumptions expected to hold then.

The evidence model is deterministic by construction. EA-11, the published execution model, reduces each governed execution to a set of component hashes composed into a Merkle-rooted evidence state, so that any party holding the public keys and the published chain can recompute and verify without trusting the operator. (The internal arithmetic is proprietary; the evidence formats and verification procedures are public.)

Figure 6-3. Public Proof Chain

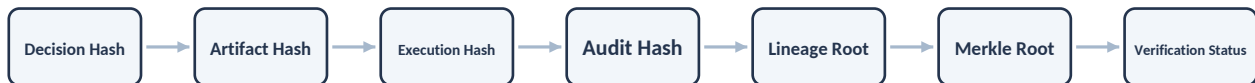


Figure 6-3. The published proof chain: component hashes compose into a Merkle-rooted, independently verifiable evidence state (EA-11).

Component	Primitive	Role
Authorization artifact signature	Ed25519	Binds the decision to the control plane's key; verifiable against published JWKS.
Audit chain hash	SHA3-512	Tamper-evident ordering of decisions and executions.
Secondary evidence chain	BLAKE2b-512	Independent second hash lineage; dual-chain cross-checking.
Post-quantum signatures	ML-DSA-65 · SLH-DSA (RFC-EG-0301)	Long-horizon verifiability of evidence beyond the classical-cryptography assumption.
Evidence state	EA-11 Merkle root	Single root summarizing the component hashes of a governed execution; supports selective disclosure (RFC-EG-0310).

Table 6-1. Execution proof components as published. Primitives are named, not implied.

6.4 Execution lineage

Data lineage, which the FSB report commends for auditability and root-cause analysis, has an execution-layer analogue. Execution lineage (RFC-EG-0010) is the tamper-evident, ordered record of governed executions:

canonical event encoding, dual-hash chains, signatures, sequence integrity, and chain-head metadata. Lineage answers, for actions, the questions data lineage answers for data: what happened, in what order, initiated by whom, authorized under what policy. For multi-agent systems it reconstructs delegation: which agent tasked which sub-agent, under which inherited authority. An open-source reference verifier validates lineage records end to end, so that verification does not depend on tooling controlled by the operator. Test vectors for authorization artifacts, lineage, and selective-disclosure bindings are published as RFC-EG-9001, RFC-EG-9010, and RFC-EG-9020.

6.5 Independent verification and auditability

The architecture's evidentiary claims are designed to be checked by parties other than the operator. Requirements for independent verification are specified in RFC-EG-0820; conformance levels in RFC-EG-0810; a certification framework in RFC-EG-0800; and a compatibility designation (Execution Governance Compatible) in RFC-EG-0700. On the operating control plane, verification is exercised in public: health, proof-viewer, and machine-readable evidence endpoints are accessible without credentials, and weekly public proof drops (the LPG-EG series) publish signed evidence on a stated cadence with deterministic verification procedures whose outcomes are PASS or FAIL. For a supervisor, the significance is procedural: examination can move from interviewing the institution about its controls to re-verifying the institution's evidence directly.

6.6 The runtime trust stack

Figure 6-4. Runtime Trust Stack

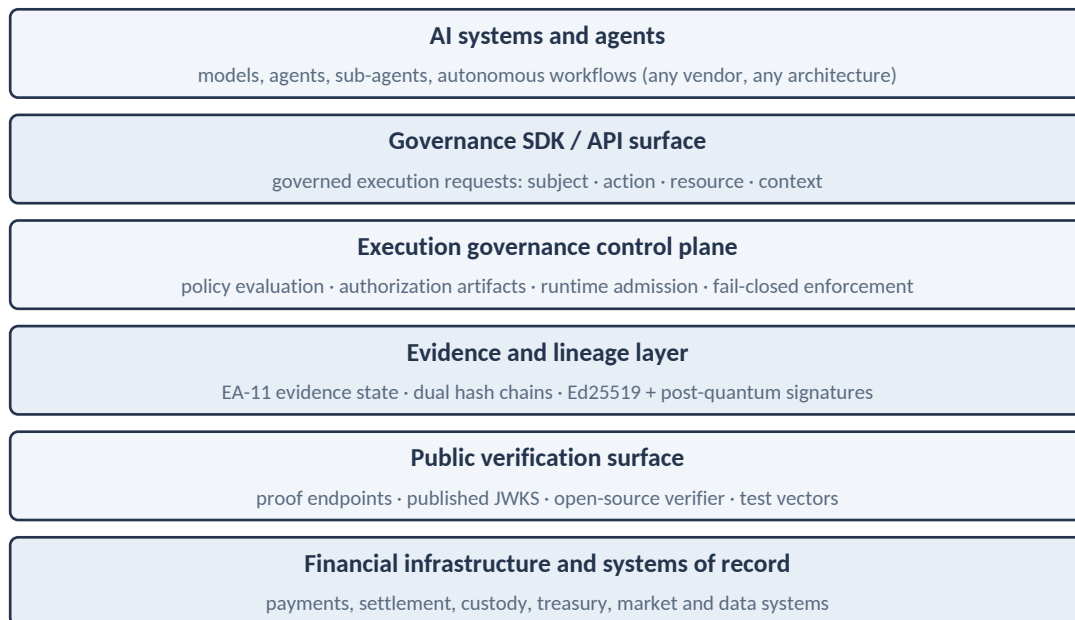


Figure 6-4. The control plane sits between AI systems and the infrastructure they act upon; evidence and verification surfaces face outward to auditors and supervisors (RFC-EG-0020, RFC-EG-1000).

The stack yields the before/during/after control set that the paper's title triad summarizes:

Phase	Controls	Governance artifacts produced
Before execution	Identity binding (RFC-EG-0120); policy evaluation; materiality and scope expressed as subject-action-resource-context; device attestation where required (RFC-EG-0220).	Governance decision; authorization artifact (RFC-EG-0100).
During execution	Runtime admission (RFC-EG-0200); fail-closed default; per-step authorization of intermediate actions; graduated degradation to human checkpoints as policy.	Admission records; denial records with policy-violation reasons.
After execution	Evidence state computation (EA-11); lineage append (RFC-EG-0010); publication to verification surfaces; selective disclosure to authorized parties (RFC-EG-0310).	Signed receipts; dual hash chains; Merkle-rooted evidence state; lineage records.

Table 6-2. Before/During/After execution controls and the governance artifacts each phase produces.

6.7 Governance decision flow

The decision flow is deliberately small. A request that satisfies the active policy set yields APPROVED and an authorization artifact. A request that violates policy yields DENIED, with a signed record carrying the policy-violation reason. A request that cannot be evaluated, because the policy set is unavailable, the identity token fails validation, the artifact has expired, or the control plane is degraded, yields FAIL_CLOSED: the system treats uncertainty as denial. A request whose policy requires a human checkpoint yields PENDING_REVIEW and waits at the boundary. There are no other terminal states, and there is no path around the control plane within the governed perimeter. The small state space matters for supervision: four outcomes, each signed, each countable, each reproducible from the evidence chain. Supervisory questions about agent conduct reduce to queries over a closed vocabulary rather than interpretation of free-form logs.

6.8 Governance artifacts

The architecture is defined as much by the artifacts it emits as by the decisions it makes. Each artifact is specified, hash-bound to its neighbors, and independently verifiable.

Artifact	Produced	Specification	Verifiable by
Governance decision record	Every request (all four outcomes)	RFC-EG-0100 · EA-11	Any party, against published keys
Authorization artifact	APPROVED decisions	RFC-EG-0100	Downstream enforcement points; auditors
Denial record with reason	DENIED / FAIL_CLOSED decisions	RFC-EG-0100 · RFC-EG-0200	Risk, audit, supervisors
Identity token	Per subject, per session or per action	RFC-EG-0120	Control plane; relying systems
Attestation record	Where device provenance is required	RFC-EG-0220	Control plane; auditors
Lineage record	Every governed execution, append-only	RFC-EG-0010	Open-source reference verifier

Artifact	Produced	Specification	Verifiable by
Evidence state (Merkle root)	Per execution; rolled into chain heads	EA-11	Any party; recomputable from components
Selective-disclosure proof	On demand, for a defined audience	RFC-EG-0310	The disclosed-to party, without full-chain access
Public proof drop	Weekly cadence (LPG-EG series)	LPG-EG-001	Public; deterministic PASS/FAIL procedure

Table 6-3. Governance artifacts: what the layer emits, under which specification, and who can verify each.

7. Financial Infrastructure Applications

The domain profile for financial infrastructure is published as RFC-EG-0610. The pattern across all applications below is identical, which is the point: one governance layer, one evidence format, one verification procedure, applied to every AI-initiated action regardless of rail or asset class. High-value and settlement-affecting actions are authorized before they execute, with per-tenant isolation and signed, verifiable evidence for every decision. (Direct payment-rail and stablecoin integrations are on the published roadmap; the descriptions below characterize the architecture's application, not deployed integrations.)

7.1 Payments

Each AI-initiated payment instruction is an authorization request: originator agent (subject), instruct-payment (action), rail and account (resource), amount, counterparty, and velocity conditions (context). Value thresholds, counterparty allowlists, and velocity limits are policy; instructions outside policy are denied before reaching the rail. This generalizes the FSB's recommended controls (dual authorization above thresholds, restricted direct agent access to payment systems) into uniform infrastructure.

7.2 Settlement

Where finality makes remediation impossible, authorization must precede execution as a matter of architecture. Pre-settlement checks become deterministic policy gates, and every released or withheld instruction carries a signed record, giving participants and overseers a shared evidentiary basis after incidents.

7.3 Custody and treasury

Movement of client assets and use of firm liquidity are bounded by policy on asset class, size, destination, and market condition. An agent rebalancing liquidity operates inside an envelope of authorized actions; attempts outside it produce signed denials that are themselves supervisory signal.

7.4 Digital assets, stablecoins, and CBDCs

Programmable money implies programmable actors, and reversibility is often absent by design. Execution governance supplies the missing institutional layer: authorization artifacts and evidence chains for AI-initiated mint, burn, transfer, and market operations, with post-quantum signature profiles (RFC-EG-0301) matched to the long horizons of monetary infrastructure. For central banks, the same evidence surfaces support oversight of AI acting within CBDC ecosystems.

7.5 Cross-border payments

Actions cross jurisdictions; evidence must too. Signed, portable execution receipts give each jurisdiction's supervisor independently verifiable records of what an agent did, without requiring access to the operator's internal systems. Federation and distributed enforcement across organizational boundaries are specified in RFC-EG-0500.

7.6 Fraud detection and prevention

The FSB's own case study shows agentic fraud detection with human-in-the-loop rule approval. Execution governance hardens the pattern: the detecting agent's proposed interventions (blocking a card, freezing a transfer) are themselves governed executions, so defensive automation cannot become an ungoverned actor

with write access to customer accounts.

7.7 Compliance, AML, and sanctions

Screening obligations become pre-execution policy: transactions to sanctioned counterparties are not flagged after the fact but denied before execution, with a signed record of the denial and its policy basis. Lineage supports the audit and lookback obligations that AML regimes impose, and selective disclosure (RFC-EG-0310) permits sharing proofs with authorities without exposing unrelated records.

7.8 Operational risk and resilience

Fail-closed defaults bound the blast radius of agent malfunction and of cyber events that manipulate agents (memory poisoning, prompt injection): a compromised agent inside the governed perimeter still cannot execute outside policy. Denial-rate telemetry provides early warning of behavioral drift, and evidence chains accelerate post-incident root-cause analysis.

Financial risk	Where it concentrates	Execution governance treatment
Irreversible loss (finality)	Settlement, payments, digital assets	Authorization strictly precedes execution; fail-closed default.
Unauthorized asset movement	Custody, treasury, payments	Resource-scoped artifacts; value and destination policy; signed denials.
Sanctions / AML breach	Cross-border, correspondent banking	Screening as pre-execution policy; provable denial records; lineage for lookback.
Runaway automation	Trading, fraud response, ops automation	Per-action admission; graduated degradation to human checkpoints; policy-set kill switch.
Evidential failure after incident	All of the above	Dual-hash, signed, Merkle-rooted evidence verifiable by third parties (EA-11, RFC-EG-0010).
Cross-institution contagion	Multi-agent, multi-firm workflows	Federated enforcement and portable proofs across boundaries (RFC-EG-0500).

Table 7-1. Financial risk mapping: where agentic execution risk concentrates and how the layer treats it.

8. Governance APIs, Implementation Model, and Open Standards

8.1 Governance APIs

The control plane exposes a small, auditable surface. A governed execution request binds subject, action, resource, and context; the response is a governance decision (APPROVED, DENIED, FAIL_CLOSED, or PENDING_REVIEW) and, when permitted, an authorization artifact. Public, credential-free surfaces exist for verification: a health endpoint reporting infrastructure status, a human-readable proof viewer, and a machine-readable public evidence endpoint whose responses verify against published JWKS keys. Signed decisions are typically returned in well under a second (published figures state typical signed-decision latency under 100 milliseconds), which is the operating regime financial workflows require. A typed JavaScript/TypeScript SDK for governed execution, artifact verification, and proof verification is in private beta.

8.2 Implementation model

Adoption does not require replacing AI systems or rails. The control plane interposes at the integration points where agents already call tools and APIs: the agent's tool-execution layer requests authorization; enforcement wraps the credentialed interfaces agents use; evidence flows to the lineage layer automatically. An institution can begin with observation (evidence generation only), proceed to enforcement on a bounded use case such as payments above a threshold, and expand policy coverage as confidence grows: the staged-approval path the FSB report commends, expressed in infrastructure. Per-tenant isolation keeps one institution's policies and evidence segregated from another's.

8.3 Open standards and conformance

The specification corpus is published as an open, DOI-registered RFC series with a chartered editorial process (RFC-EG-0001, 0002) and errata and lifecycle procedures (RFC-EG-0003). Conformance is specified in graded levels (RFC-EG-0810) under a certification framework (RFC-EG-0800) with independent verification requirements (RFC-EG-0820); machine-readable conformance profiles (EG-CORE, and domain profiles including financial infrastructure) are published in open source with mappings to recognized frameworks including the NIST AI Risk Management Framework. Published guidance is explicit that certification may not be claimed unless formal certification has been completed. The intent of the open publication posture is that the architecture can be evaluated, criticized, implemented, and tested by parties with no relationship to its author, which is the property standard-setters should demand of any governance layer proposed for systemic infrastructure.

8.4 Operational assurance

The assurance model can be summarized as a matrix of the questions an institution, auditor, or supervisor needs answered, the mechanism that answers each, and the party able to verify the answer without trusting the operator.

Assurance question	Mechanism	Independent verification
Was this action authorized before it ran?	Authorization artifact bound to subject, action, resource, context (RFC-EG-0100).	Signature check against published JWKS; test vectors RFC-EG-9001.
Could the agent have acted without authorization?	Runtime admission; fail-closed default (RFC-EG-0200).	Denial records on the public proof surface; both ALLOW and DENY are signed.
What did the agent actually do, in order?	Execution lineage: canonical encoding, dual hash chains, sequence integrity (RFC-EG-0010).	Open-source reference verifier; test vectors RFC-EG-9010.
Has the record been altered since?	SHA3-512 and BLAKE2b-512 dual chains; EA-11 Merkle-rooted evidence state.	Recomputation from components by any party.
Will the evidence still verify in 15 years?	Post-quantum signature profile: ML-DSA-65, SLH-DSA (RFC-EG-0301).	Verification against the published profile; DOC-EG-005 states the rationale.
Can proofs be shared without exposing everything?	Selective-disclosure proof format (RFC-EG-0310).	Disclosed-to party verifies the binding; test vectors RFC-EG-9020.
Is the control plane itself operating?	Public health endpoint; weekly public proof drops (LPG-EG series).	Anyone, without credentials, on a stated cadence.
Does an implementation conform?	Conformance levels (RFC-EG-0810); independent verification requirements (RFC-EG-0820).	Third-party assessment under the published framework (RFC-EG-0800).

Table 8-1. Operational assurance matrix: each assurance question is answered by a published mechanism and is verifiable by a party other than the operator.

9. Mapping to the FSB Sound Practices

This section maps each of the twelve sound practices to the corresponding execution governance doctrine, in the structure requested by supervisory readers: current practice as described in the consultation report, the execution governance enhancement, the relevant published doctrine, an implementation example, and the operational outcome. The mapping presents execution governance as an operationalization of the sound practices for agentic AI, not as a substitute for any of them.

9.1 Sound Practice 1: Strategic direction and oversight

Current practice	Board and senior management align AI adoption with business model, risk appetite, and strategy, including articulating prohibited AI use cases and considering the level of autonomy, materiality, and customer impact.
EG enhancement	Risk appetite is compiled into machine-readable policy enforced on every execution. Prohibited use cases become classes of action that cannot execute within the governed perimeter; boundary changes take effect at the control plane without re-engineering agents.
Relevant doctrine	EG-001 (category definition); DOC-EG Architecture (10.5281/zenodo.20252295)
Implementation example	A board prohibition on fully automated decision-making in a critical business area is expressed as a policy requiring PENDING_REVIEW (human checkpoint) for that action class; the agent can propose but cannot execute.
Operational outcome	Risk appetite statements become testable controls; the board can be shown evidence, not assurances, that boundaries hold.

Section 9.1. Mapping of Sound Practice 1 to published execution governance doctrine.

9.2 Sound Practice 2: Governance and accountability

Current practice	Clear roles and responsibilities on existing structures such as the three lines of defence; audit, risk, and testing remain independent from development and empowered to challenge.
EG enhancement	The control plane is a structural embodiment of independence: authorization is decided and recorded outside the systems being governed. Second-line policy owners control the policy set; the first line cannot alter evidence; audit verifies records cryptographically rather than by inquiry.
Relevant doctrine	RFC-EG-0020 (architectural decomposition of the control plane)
Implementation example	Risk management owns policy administration in the governance console; internal audit re-verifies lineage with the open-source verifier against published keys.
Operational outcome	Independent challenge is enforced by architecture rather than by organizational discipline alone.

Section 9.2. Mapping of Sound Practice 2 to published execution governance doctrine.

9.3 Sound Practice 3: Incorporation of AI risks into risk management framework

Current practice	AI risks incorporated into frameworks; inventories and enhanced documentation and logging for GenAI and agentic use cases, including versioning, rollback, and chain-of-thought logging; staged approvals.
EG enhancement	Documentation becomes a by-product of operation: every governed execution generates its own signed record, including denials with policy-violation reasons. Agent inventories gain runtime ground truth: what each agent did, was denied, and under which policy version.
Relevant doctrine	RFC-EG-0010 (execution lineage); DOC-EG-003 (agentic doctrine)
Implementation example	The report's agent 'certification' concept (approval within defined boundaries, recorded as inventory attributes) maps to an authorization scope: certified boundaries are the agent's policy envelope, and certification status is checkable against live enforcement records.
Operational outcome	Inventories reflect observed behavior rather than declared intent; staged approvals are enforced stages.

Section 9.3. Mapping of Sound Practice 3 to published execution governance doctrine.

9.4 Sound Practice 4: Organisational adaptability

Current practice	Institutions learn and adapt oversight, governance, and risk practices as AI evolves, including governance mechanisms that adapt to new forms of AI.
EG enhancement	The governed boundary is the invariant: new model generations, frameworks, or agent architectures are onboarded by policy at the control plane, without rebuilding controls per technology generation.
Relevant doctrine	DOC-EG-FRAMEWORK (national infrastructure framework)
Implementation example	A new agentic framework is admitted first in observation mode (evidence only), then under narrow policy, then at target scope: adaptation as configuration, with the same evidence format throughout.
Operational outcome	Control architecture survives technology churn; adaptation cost falls from engineering to policy.

Section 9.4. Mapping of Sound Practice 4 to published execution governance doctrine.

9.5 Sound Practice 5: Materiality and risk assessment

Current practice	Systematic assessment of materiality and risk at inception and thereafter, considering autonomy, capabilities, complexity, access to systems and data, and scope of action.
EG enhancement	The assessment inputs the FSB names are first-class, machine-readable quantities: an agent's scope of action is its authorized subject-action-resource-context set, and its actual exercised scope is measurable from lineage. Reassessment triggers can key on observed drift between the two.
Relevant doctrine	RFC-EG-0100 (authorization artifact); EA-11 (deterministic evidence state)
Implementation example	The report's example of a low-materiality agent with high-risk access is directly addressed: access exists only as explicit authorized tuples, so 'what could this agent reach' is a query, not an investigation.
Operational outcome	Materiality and risk assessments rest on measured scope rather than estimated scope.

Section 9.5. Mapping of Sound Practice 5 to published execution governance doctrine.

9.6 Sound Practice 6: Selection

Current practice	Selection weighs objectives, materiality, risks, build-versus-buy, and recognizes that simpler, more explainable systems may suffice.
EG enhancement	Selection criteria gain a governability dimension: can the candidate operate under external authorization, runtime admission, and evidence generation? Conformance designations give procurement a concrete artifact to require.
Relevant doctrine	RFC-EG-0700 (Execution Governance Compatible); RFC-EG-0810 (conformance levels)
Implementation example	An RFP for an agentic system requires operation under fail-closed external authorization and lineage emission, verified against published test vectors (RFC-EG-9001/9010) during evaluation.
Operational outcome	Governability is assessed before commitment, when leverage is highest, rather than retrofitted.

Section 9.6. Mapping of Sound Practice 6 to published execution governance doctrine.

9.7 Sound Practice 7: Data governance

Current practice	Fit-for-purpose data across the lifecycle; data lineage for auditability and root cause; policy-aware data whose metadata can act as guardrails; attention to how agents ingest, memorise, and re-use data.
EG enhancement	Execution lineage extends data lineage to the action layer: which agent, under which authorization, read or wrote which data resource. Policy-aware data acquires an enforcement point: sensitivity and permitted-use metadata are honored because access is mediated per execution.
Relevant doctrine	RFC-EG-0010 (lineage); RFC-EG-0310 (selective-disclosure proof format)
Implementation example	An agent's attempt to query customer records broader than its task requires (the report's scope-creep example) is denied by resource-scoped policy and logged as a signed denial for review.
Operational outcome	Data governance obtains runtime teeth; scope creep surfaces as evidence rather than as forensics.

Section 9.7. Mapping of Sound Practice 7 to published execution governance doctrine.

9.8 Sound Practice 8: Explainability and transparency

Current practice	Explainability as a continuum with compensating controls where limited; for agentic systems, tracing intermediate decisions and steps; transparency tailored to stakeholders.
EG enhancement	Where model internals resist explanation, the governed action record explains conduct: the full sequence of authorized and denied actions, with policy bases, reconstructs the execution path (the report's local explainability for agents). This is precisely a compensating control in the report's sense, strongest where model explainability is weakest.
Relevant doctrine	DOC-EG-003; RFC-EG-0010
Implementation example	For a contested agent outcome, the institution produces the signed sequence of tool invocations and intermediate decisions that led to it, supporting contestability and redress channels.
Operational outcome	Conduct-level explainability holds even for opaque models; stakeholder transparency is backed by evidence.

Section 9.8. Mapping of Sound Practice 8 to published execution governance doctrine.

9.9 Sound Practice 9: Performance management

Current practice	Performance evaluated proportionately, with ante hoc thresholds preferred; for agents, assessment must consider whether the actions taken were appropriate, not only whether objectives were achieved; routinised assessment of specific actions.
EG enhancement	Per-action authorization records are the measurable substrate for action-level assessment: rates of denial, near-boundary actions, and policy-exception requests are computable, continuous performance signals with ante hoc semantics (the threshold is the policy).
Relevant doctrine	EA-11 (deterministic evidence state); LPG-EG series (published verification cadence)
Implementation example	Routinised assessment becomes a query over lineage: all actions by agent class X against policy version Y, with deterministic PASS/FAIL verification procedures over the evidence itself.
Operational outcome	The report's 'actions, not just outputs' requirement becomes operational, at agent scale.

Section 9.9. Mapping of Sound Practice 9 to published execution governance doctrine.

9.10 Sound Practice 10: Human oversight

Current practice	Meaningful oversight matched to materiality and autonomy; forms from human-in-the-loop to human-in-command, kill switch, contestability; boundaries and approval checkpoints for agents; controls on agents executing financial transactions; audit trails; treating agents as synthetic employees.
EG enhancement	Oversight is re-anchored where it scales: humans govern the policy set, approve exceptions at defined checkpoints (PENDING_REVIEW), and audit evidence; the control plane applies policy uniformly to every execution. Kill-switch capability is the resting state of a fail-closed system. The synthetic-employee framing maps naturally: identity tokens, scoped authority, and a personnel file of signed conduct records.
Relevant doctrine	DOC-EG-003; RFC-EG-0200 (runtime admission); RFC-EG-0120 (identity tokens)
Implementation example	Dual authorization above a value threshold is policy: the artifact issues only upon a second approval; below threshold, execution is autonomous and evidenced. Rubber-stamping analysis runs over recorded approval patterns.
Operational outcome	Accountability remains human, as the report requires, while per-decision review ceases to be the bottleneck.

Section 9.10. Mapping of Sound Practice 10 to published execution governance doctrine.

9.11 Sound Practice 11: Cyber and ICT risk management

Current practice	AI-specific cyber practice: dynamic IAM for agents, least privilege for agents and sub-agents, monitoring and logging, controlled environments to observe agent behavior pre-production, information sharing.
EG enhancement	Least privilege and dynamic IAM are enforced per execution rather than per session: authority is granted action by action, against live context. A compromised, poisoned, or prompt-injected agent inside the perimeter still cannot execute outside policy, because enforcement is external to the agent; its attempts become signed denials, which are detection signal.
Relevant doctrine	RFC-EG-0120 (identity token profile); RFC-EG-0220 (device attestation chain)
Implementation example	The report's controlled-environment practice maps to observation-mode admission: agents run under full evidence generation with narrow enforcement before production scope is granted.
Operational outcome	The blast radius of agent compromise is bounded by policy, not by the attacker's imagination.

Section 9.11. Mapping of Sound Practice 11 to published execution governance doctrine.

9.12 Sound Practice 12: Third-party AI risk management

Current practice	Due diligence on providers; attention to performance, transparency, data quality, supply chain, concentration, and continuity; contractual arrangements where vendor control limits rollback or configuration access.
EG enhancement	Third-party and vendor-hosted agents operate under the institution's own control plane, so governance does not depend on the vendor's internal practices: the institution's policies gate execution, and the institution holds its own signed evidence. Federated enforcement extends the model across organizational boundaries, and portable proofs support due diligence, supervision, and exit.
Relevant doctrine	RFC-EG-0500 (federation and distributed enforcement); RFC-EG-0820 (independent verification requirements)
Implementation example	A vendor model update (the report's transparency concern) cannot silently expand an agent's effective authority: scope is fixed by the institution's policy, and any behavioral shift appears as changed denial patterns in evidence.
Operational outcome	Institutions retain governance of execution even where they do not control the model; concentration risk is partially offset by evidence portability.

Section 9.12. Mapping of Sound Practice 12 to published execution governance doctrine.

10. Recommendations

For standard-setters and supervisors. (1) Recognize runtime execution controls, comprising pre-execution authorization, fail-closed enforcement, and verifiable execution evidence, as a named control class for agentic AI in the final report, complementing lifecycle governance. (2) State a fail-closed expectation for material agentic use cases touching payments, settlement, custody, or customer funds: absent explicit authorization, the action does not execute. (3) Encourage cryptographic, third-party-verifiable execution evidence for material use cases, mechanism-neutral, and note post-quantum signature profiles for evidence with long retention horizons. (4) Frame human oversight of agents as governance of the policy layer with exception checkpoints, preserving human accountability while abandoning the untenable premise of per-decision review. (5) Encourage open specifications, conformance criteria, test vectors, and independent verification for any execution governance mechanism, from any vendor, proposed for systemic infrastructure.

For financial institutions. (1) Inventory the execution boundaries where AI-initiated actions take effect, and assess which currently have an independent authorization point. (2) For agentic pilots, adopt staged admission: observation (evidence only), bounded enforcement, then scaled scope. (3) Express risk appetite, prohibited use cases, and transaction controls as machine-readable policy wherever an execution governance layer exists to enforce them. (4) Require governability in AI procurement: the ability of a candidate system to operate under external authorization and to emit verifiable evidence. (5) Treat denial records as first-class risk telemetry.

11. Conclusion

The FSB consultation report identifies, accurately, the properties of agentic AI that strain current governance: autonomy that produces unauthorized and erroneous actions, oversight that cannot scale to per-decision review, actions whose effects cannot be remediated, and evidence that depends on the governed system's own logs. Each of these is an execution-layer problem, and each has an execution-layer answer that is specifiable and testable today: authorization before execution, enforcement during execution, and cryptographic proof after execution.

This paper has presented one published, publicly verifiable embodiment of that layer. The doctrine and specifications are DOI-registered and openly licensed; the verification tooling is open source; the control plane's signed decisions are exposed for public verification; and the mapping to the FSB's twelve sound practices is, we hope, concrete enough to be adopted, contested, or improved. The underlying claim does not depend on any vendor, including this one: as financial institutions delegate execution to autonomous systems, governance of execution itself becomes part of the safety architecture of the financial system. We offer this work to the FSB and its members as a contribution to that architecture.

12. References

- Financial Stability Board (2026). Sound Practices for Responsible Adoption of Artificial Intelligence (AI): Consultation report. 10 June 2026. <https://www.fsb.org/uploads/P100626.pdf>
- Financial Stability Board (2024). The Financial Stability Implications of Artificial Intelligence. November 2024.
- Financial Stability Board (2017). Artificial Intelligence and Machine Learning in Financial Services. November 2017.
- National Institute of Standards and Technology (2023). AI Risk Management Framework (AI RMF 1.0).
- 11/11 AI Research Division (2026). 11/11 Execution Governance Architecture: A Fail-Closed Execution Governance Layer for Regulated AI and Compute Infrastructure. DOI 10.5281/zenodo.20252295.

11/11 AI Research Division (2026). 11/11 AI: A Deterministic Execution Authority Architecture for Governed AI, Data, and Autonomous Systems. DOI 10.5281/zenodo.20252639.

11/11 AI Research Division (2026). EG-001: Execution Governance: A Proposed Infrastructure Category for the Pre-Execution Authorization of Autonomous Systems. DOI 10.5281/zenodo.20453136.

11/11 AI Research Division (2026). DOC-EG-FRAMEWORK: Execution Governance: A National Infrastructure Framework for Autonomous Systems, AI Runtime Control, and Governed Execution. DOI 10.5281/zenodo.20385909.

11/11 AI Research Division (2026). DOC-EG-003: Execution Governance for Agentic AI Systems: A Doctrine for Inter-Agent Authorization, Lineage, and Audit. DOI 10.5281/zenodo.20261800.

11/11 AI Research Division (2026). DOC-EG-004: Public Proof Endpoints. DOI 10.5281/zenodo.20348548.

11/11 AI Research Division (2026). DOC-EG-005: The Cryptographic Half-Life of AI: Post-Quantum Authentication and Long-Term Trustworthiness of AI Artifacts. DOI 10.5281/zenodo.21053032.

11/11 AI Research Division (2026). EA-11: Execution Arithmetic: A Governance-Native Execution Model for Deterministic Authorized Computation. DOI 10.5281/zenodo.20418630.

11/11 AI Research Division (2026). RFC-EG-0001: Charter of the RFC-EG Series. DOI 10.5281/zenodo.20260882.

11/11 AI Research Division (2026). RFC-EG-0002: Editorial Process. DOI 10.5281/zenodo.20348799.

11/11 AI Research Division (2026). RFC-EG-0003: Errata and Lifecycle Procedures. DOI 10.5281/zenodo.20350950.

11/11 AI Research Division (2026). RFC-EG-0010: Execution Lineage Specification. DOI 10.5281/zenodo.20264726.

11/11 AI Research Division (2026). RFC-EG-0020: Architectural Decomposition of the Control Plane. DOI 10.5281/zenodo.20350744.

11/11 AI Research Division (2026). RFC-EG-0100: Authorization Artifact Specification. DOI 10.5281/zenodo.20266345.

11/11 AI Research Division (2026). RFC-EG-0120: Identity Token Profile for Execution Governance. DOI 10.5281/zenodo.20348838.

11/11 AI Research Division (2026). RFC-EG-0200: Runtime Admission Specification. DOI 10.5281/zenodo.20275415.

11/11 AI Research Division (2026). RFC-EG-0220: Device Attestation Chain. DOI 10.5281/zenodo.20351100.

11/11 AI Research Division (2026). RFC-EG-0301: Post-Quantum Signature Profile. DOI 10.5281/zenodo.20836771.

11/11 AI Research Division (2026). RFC-EG-0310: Selective-Disclosure Proof Format. DOI 10.5281/zenodo.20348760.

11/11 AI Research Division (2026). RFC-EG-0500: Federation and Distributed Enforcement. DOI 10.5281/zenodo.20278155.

11/11 AI Research Division (2026). RFC-EG-0610: Domain Profile: Financial Infrastructure. DOI 10.5281/zenodo.20314865.

11/11 AI Research Division (2026). RFC-EG-0700: Execution Governance Compatible (EGC). DOI 10.5281/zenodo.20467986.

11/11 AI Research Division (2026). RFC-EG-0800: Execution Governance Certification Framework. DOI 10.5281/zenodo.20468025.

11/11 AI Research Division (2026). RFC-EG-0810: Execution Governance Conformance Levels. DOI 10.5281/zenodo.20468056.

11/11 AI Research Division (2026). RFC-EG-0820: Execution Governance Independent Verification Requirements. DOI 10.5281/zenodo.20468085.

11/11 AI Research Division (2026). RFC-EG-1000: Execution Governance Reference Architecture. DOI 10.5281/zenodo.20468141.

11/11 AI Research Division (2026). RFC-EG-9001 / 9010 / 9020: Test Vectors for Authorization Artifacts, Execution Lineage, and Selective-Disclosure Bindings. DOIs 10.5281/zenodo.20356881, 10.5281/zenodo.20356909, 10.5281/zenodo.20356979.

11/11 AI Research Division (2026). LPG-EG-001: Public Proof Drop Cadence and Public Proof Drops #001-#006. DOIs 10.5281/zenodo.20434149 through 10.5281/zenodo.21079616 (see Technical Evidence Appendix).

11/11 AI Research Division (2026). IDX-EG-001: Doctrine Index v1.11, master catalog of 40 canonical artifacts. DOI 10.5281/zenodo.21078499.

11/11 AI (2026). Public verification surfaces: control.11aiblockchain.com (health, proof viewer, public evidence endpoint); reference verifier and conformance profiles at github.com/AtlasQuantumProtocol.

License and Marks

Execution Governance™, Governed Execution™, EA-11™, Execution Arithmetic™, and Execution Evidence State™ are trademarks of 11 AI AND BLOCKCHAIN DEVELOPMENT LLC, Patent Pending. Intellectual property protected by the 11 AI Blockchain Developments Land and IP Trust.

Public Technical Evidence Appendix

Published doctrine, specifications, proof endpoints, and repositories supporting the 11/11 AI response to the FSB consultation on Sound Practices for Responsible Adoption of Artificial Intelligence (AI)

11 AI AND BLOCKCHAIN DEVELOPMENT LLC

11/11 AI Research Division

All items in this appendix are public: DOI-registered, open-source, or accessible without credentials.

30 N Gould St Ste R, Sheridan, WY 82801 • quantum@11aiblockchain.com • 11aiblockchain.com

July 2026

Execution Governance™ • Governed Execution™ • EA-11™ • Execution Arithmetic™ • Execution Evidence State™ • Patent Pending

A. Identification

Document	Public Technical Evidence Appendix to the 11/11 AI FSB consultation submission.
Purpose	Enable reviewers to verify every capability claim in the consultation response and technical white paper against public sources, without access to any private system.
Publisher	11 AI AND BLOCKCHAIN DEVELOPMENT LLC · 11/11 AI Research Division
Canonical archive	Zenodo community 11-11-ai: zenodo.org/communities/11-11-ai/records (52+ DOI records). Machine-readable listing: zenodo.org/api/communities/11-11-ai/records
Master catalog	IDX-EG-001 Doctrine Index v1.11, 40 canonical artifacts. DOI 10.5281/zenodo.21078499
Licensing	Records are openly licensed (CC BY / CC BY-ND as marked on each record).

Section A. Identification of this appendix and the canonical public archive.

B. Foundational Doctrine

Code	Title	DOI
DOC-EG	11/11 Execution Governance Architecture: A Fail-Closed Execution Governance Layer for Regulated AI and Compute Infrastructure	10.5281/zenodo.20252295
DOC-EG	11/11 AI: A Deterministic Execution Authority Architecture for Governed AI, Data, and Autonomous Systems	10.5281/zenodo.20252639
EG-001	Execution Governance: A Proposed Infrastructure Category for the Pre-Execution Authorization of Autonomous Systems	10.5281/zenodo.20453136
DOC-EG-FRAMEWORK	Execution Governance: A National Infrastructure Framework for Autonomous Systems, AI Runtime Control, and Governed Execution	10.5281/zenodo.20385909
DOC-EG-003	Execution Governance for Agentic AI Systems: A Doctrine for Inter-Agent Authorization, Lineage, and Audit	10.5281/zenodo.20261800
DOC-EG-004	Public Proof Endpoints	10.5281/zenodo.20348548
DOC-EG-005	The Cryptographic Half-Life of AI: Post-Quantum Authentication and Long-Term Trustworthiness of AI Artifacts	10.5281/zenodo.21053032
EA-11	Execution Arithmetic: A Governance-Native Execution Model for Deterministic Authorized Computation	10.5281/zenodo.20418630
OPN-EG-001	Open Framework Response for Execution Governance Infrastructure	10.5281/zenodo.20684336

Section B. Foundational doctrine and architecture records.

C. RFC-EG Specification Series

Process and charter

Code	Title	DOI
RFC-EG-0001	Charter of the RFC-EG Series	10.5281/zenodo.20260882
RFC-EG-0002	Editorial Process	10.5281/zenodo.20348799
RFC-EG-0003	Errata and Lifecycle Procedures	10.5281/zenodo.20350950

Core specifications

Code	Title	DOI
RFC-EG-0010	Execution Lineage Specification	10.5281/zenodo.20264726
RFC-EG-0020	Architectural Decomposition of the Control Plane	10.5281/zenodo.20350744
RFC-EG-0100	Authorization Artifact Specification	10.5281/zenodo.20266345
RFC-EG-0120	Identity Token Profile for Execution Governance	10.5281/zenodo.20348838
RFC-EG-0200	Runtime Admission Specification	10.5281/zenodo.20275415
RFC-EG-0220	Device Attestation Chain	10.5281/zenodo.20351100
RFC-EG-0300	512 Cypher Specification	10.5281/zenodo.20277892
RFC-EG-0301	Post-Quantum Signature Profile	10.5281/zenodo.20836771
RFC-EG-0310	Selective-Disclosure Proof Format	10.5281/zenodo.20348760
RFC-EG-0400	11/11 Lang: Syntax and Semantics	10.5281/zenodo.20278038
RFC-EG-0500	Federation and Distributed Enforcement	10.5281/zenodo.20278155
RFC-EG-0510	Chain-Head Substrate Selection and Minimum Durability	10.5281/zenodo.20351133

Domain profiles

Code	Title	DOI
RFC-EG-0600	Domain Profile: Regulated AI Inference	10.5281/zenodo.20280530
RFC-EG-0610	Domain Profile: Financial Infrastructure	10.5281/zenodo.20314865
RFC-EG-0620	Domain Profile: HIPAA-Compliant Clinical AI	10.5281/zenodo.20315104
RFC-EG-0630	Domain Profile: Defense and JADC2	10.5281/zenodo.20319267

Certification and architecture

Code	Title	DOI
RFC-EG-0700	Execution Governance Compatible (EGC)	10.5281/zenodo.20467986
RFC-EG-0800	Execution Governance Certification Framework	10.5281/zenodo.20468025
RFC-EG-0810	Execution Governance Conformance Levels	10.5281/zenodo.20468056
RFC-EG-0820	Execution Governance Independent Verification Requirements	10.5281/zenodo.20468085

Code	Title	DOI
RFC-EG-0830	Execution Governance Federal Procurement Mapping	10.5281/zenodo.20468106
RFC-EG-1000	Execution Governance Reference Architecture	10.5281/zenodo.20468141

Test vectors

Code	Title	DOI
RFC-EG-9001	Test Vectors for Authorization Artifacts	10.5281/zenodo.20356881
RFC-EG-9010	Test Vectors for Execution Lineage	10.5281/zenodo.20356909
RFC-EG-9020	Test Vectors for Selective-Disclosure Bindings	10.5281/zenodo.20356979

Section C. The RFC-EG series: process, core specifications, domain profiles, certification framework, and machine-readable test vectors.

D. Live Proof Governance (LPG-EG) and Doctrine Index

Code	Title	DOI
LPG-EG-001	Public Proof Drop Cadence	10.5281/zenodo.20434149
LPG-EG-001	Public Proof Drop #001	10.5281/zenodo.20434177
LPG-EG-001	Public Proof Drop #002	10.5281/zenodo.20533425
LPG-EG-001	Public Proof Drop #003 (ISO Week 2026-W24)	10.5281/zenodo.20647625
LPG-EG-001	Public Proof Drop #004 (ISO Week 2026-W25)	10.5281/zenodo.20738945
LPG-EG-001	Public Proof Drop #005 (ISO Week 2026-W26)	10.5281/zenodo.20947022
LPG-EG-001	Public Proof Drop #006 (ISO Week 2026-W27)	10.5281/zenodo.21079616
IDX-EG-001	Doctrine Index v1.11 (latest; 40 canonical artifacts; earlier versions remain archived under their own DOIs)	10.5281/zenodo.21078499

Section D. Weekly public proof drops publish signed evidence on a stated cadence with deterministic verification procedures.

E. Public Infrastructure Surfaces

Surface	URL	Access
Main site and doctrine hub	11aiblockchain.com	Public
Developer documentation	11aiblockchain.com/docs	Public
Research archive (DOI catalog)	11aiblockchain.com/research	Public
Control plane health	control.11aiblockchain.com/health	Public, no credentials
Public proof viewer (human-readable)	control.11aiblockchain.com/proof	Public, no credentials

Surface	URL	Access
Public evidence endpoint (machine-readable)	control.11aiblockchain.com/v1/public/evidence	Public, no credentials; signed
Governance console	control.11aiblockchain.com/console	Public surface
Live demo / governed runtime	control.11aiblockchain.com/demo · control.11aiblockchain.com/governed-ai	Public
Authenticated execution gateway	POST control.11aiblockchain.com (governed execution requests)	API key required; SDK in private beta

Section E. Public verification surfaces. Signed decision records verify against the control plane's published JWKS keys.

F. Open-Source Repositories

Repository	Contents
github.com/AtlasQuantumProtocol/11-11-li-neage-verifier	Reference implementation validating canonical event encoding, dual-hash chains, Ed25519 signatures, sequence integrity, and chain-head metadata for governed execution records (RFC-EG-0010).
github.com/AtlasQuantumProtocol/11-11-governance-profiles	Machine-readable conformance profiles EG-CORE-1, EG-DEFENSE-1, EG-HEALTH-1, EG-FEDERAL-1, with requirement mappings to recognized frameworks including the NIST AI RMF.
github.com/AtlasQuantumProtocol/11-11-execution-governance-docs	Execution governance documentation.
github.com/AtlasQuantumProtocol/11ai-developer-docs	Developer documentation for the governance API surface and SDK.
github.com/AtlasQuantumProtocol/11-11-lang-EXP	Published interfaces and guardrail documentation for 11/11 Lang (implementation not published).

Section F. Public repositories under github.com/AtlasQuantumProtocol.

G. Published Cryptographic Profile

Layer	Primitive(s)	Published in
Authorization artifact signatures	Ed25519	Proof documentation; DOC-EG-004
Audit chain hashing	SHA3-512	Proof documentation; RFC-EG-0010
Secondary evidence chaining	BLAKE2b-512	Proof documentation; RFC-EG-0010
Post-quantum signatures	ML-DSA-65 · SLH-DSA (SPHINCS+ family)	RFC-EG-0301; live on the public evidence endpoint
Evidence state	EA-11 Merkle-rooted component hashes	EA-11 (10.5281/zenodo.20418630)
Enforcement default	Fail-closed (default deny)	DOC-EG Architecture; RFC-EG-0200

Section G. The cryptography is named at every layer; nothing is implied.

H. Public Verification Procedure

INPUT: a workstation with HTTPS access. No credentials, accounts, or agreements are required.

1. Retrieve `control.11aiblockchain.com/health`. Outcome: a status response indicating the control plane is operating. PASS if reachable and healthy; otherwise FAIL.
2. Retrieve the machine-readable public evidence endpoint, `control.11aiblockchain.com/v1/public/evidence`. Outcome: a signed evidence document including signature families Ed25519, ML-DSA, and SLH-DSA and an overall verification status. PASS if `verification_status` reads VERIFIED and each signature family reads VALID; otherwise FAIL.
3. Retrieve the published JWKS keys from the control plane and verify the signatures in step 2 against them using standard tooling or the open-source reference verifier. PASS on signature validity; otherwise FAIL.
4. Open the public proof viewer, `control.11aiblockchain.com/proof`, and confirm that individual decision records (both APPROVED and denial outcomes) display decision hashes, artifact hashes, and lineage roots consistent with step 2. PASS on consistency; otherwise FAIL.
5. Cross-check any DOI cited in the submission against the Zenodo community listing (zenodo.org/api/communities/11-11-ai/records). PASS if the record resolves and matches the cited title; otherwise FAIL.

On any FAIL outcome, file per RFC-EG-0003 (Errata and Lifecycle Procedures) to quantum@11aiblockchain.com.

Section H. Deterministic verification procedure. Every claim in the submission is checkable by this route or by the cited DOI records.

I. License and Marks

Execution Governance™, Governed Execution™, EA-11™, Execution Arithmetic™, and Execution Evidence State™ are trademarks of 11 AI AND BLOCKCHAIN DEVELOPMENT LLC, Patent Pending. Intellectual property protected by the 11 AI Blockchain Developments Land and IP Trust.

Zenodo records are openly licensed as marked on each record (CC BY 4.0 or CC BY-ND 4.0). Open-source repositories carry their own license files. Nothing in this appendix asserts third-party certification, regulatory endorsement, or market adoption.

Section I. Licensing and trademark statement.