

Governing AI access to systems of record

The Omnafy security and governance architecture

Audience	Security and compliance leaders
Version	1.0
Published	August 31, 2026
Online	omnafy.com/whitepapers/governed-ai-access/

Abstract

A new employee gets an account, a role, a manager, and a paper trail before they touch a system of record. An AI caller gets an API key and a hope. That asymmetry is the reason most organizations that want AI working in their real business systems cannot allow it, and it is not a model problem or an integration problem. It is a governance problem: AI callers arrive with no identity, no permissions, no approval path, and no audit trail.

This paper describes the architecture Omnafy uses to close that gap. Every AI caller in an organization, whether a staff member's desktop AI client or an autonomous agent, becomes an authenticated principal behind a single gateway. Every tool that principal can even see carries a risk tier ratified by a human. Consequential actions are proposed and approved before they happen. Judgment calls stay permanently human by construction rather than by policy. And every action, allowed or denied or queued, lands in an append-only record naming who acted, on whose behalf, and who approved it.

The paper is written for the person who has to decide whether AI touches production systems. It states the five invariants the whole design is checked against, draws the trust boundaries and enumerates exactly what crosses them, describes the identity and authorization mechanisms in enough detail to argue with, summarizes the ten threats we designed against, and publishes the risks we accept rather than mitigate. The last of those is the section we would read first if we were you.

1. The problem, stated precisely

The usual framing is that AI agents are risky because models are unpredictable. That framing leads to the wrong controls. Model unpredictability is real, but an unpredictable actor with no authority can do nothing, and a perfectly predictable actor with unbounded authority can do anything. The risk lives in the authority, not in the prediction.

What actually happens inside organizations today looks like this. A staff member connects a desktop AI client to an internal database with a connection string. An engineering team stands up an MCP server in front of a ticketing system so an agent can triage. A vendor integration gets an API key with the permissions that were convenient at setup time. None of these callers has an identity that maps to a person. None has permissions narrower than the credential it holds. None produces a record that survives the session. When something goes wrong in a system of record, the honest answer to "who did this" is that a service account did it, which is true of every action and therefore evidence of nothing.

Four questions decide whether AI access is defensible, and none of them is answered by a demo:

Who exactly is this caller? What exactly can it do? Who approved that? Where is the record?

Omnafy exists to make those four questions answerable mechanically, for every AI caller in the organization, before the incident rather than after it.

2. Design principles: five invariants

The security posture of the platform hangs off five named guarantees. They are numbered, they are short, and every plane, contract, and design decision is checked against them. A design that violates one is wrong even when it is convenient.

ID	INVARIANT	FORMAL STATEMENT
I1	Every caller is a principal	Every call executes under exactly one authenticated principal. Humans, agents, and services share one identity model, and no anonymous or ambient-authority path exists, by construction.
I2	No side doors	Every tool interaction, including an interactive staff session, an autonomous agent, and our own administrative API, traverses the gateway and both of its enforcement points. No route reaches a tool server directly.
I3	Non-escalation	For every agent-created definition <code>d</code> with creator <code>p</code> , <code>effective_scope(d) ⊆ effective_scope(p)</code> , verified by computable scope comparison at draft time and re-verified at approval.
I4	Total, attributable audit	Every action, allowed or denied or queued, emits an append-only audit event carrying the full principal chain: caller, on-behalf-of, approver.
I5	Regulated data stays in the customer's boundary	Regulated data never enters the Omnafy control plane. Only redacted metadata crosses the deployment boundary, and model access occurs only through the customer's own contractual path.

The reasoning behind each is short enough to state in a sentence. I1 comes first because attribution is the precondition for everything else: a call that cannot be attributed cannot be scoped, approved, or audited. I2 exists because an enforcement point that can be routed around is not an enforcement point. I3 exists because delegation must never mint authority, and without it, agents that draft agents form a privilege escalation ladder. I4 exists because a trail with gaps fails exactly when something has gone wrong, and attribution is what makes a record answerable rather than merely complete. I5 exists because a governance layer must not itself become the exfiltration channel it was installed to prevent.

Two enforcement rules travel with the invariants and are inherited by every component.

Fail closed. When policy cannot be evaluated, the answer is deny, and the denial is audited. Replicated policy state unavailable, an unknown tool, a manifest without a ratified tier, a scope comparison the evaluator cannot compute: each produces a denial, never a default allowance. Allowance always requires an affirmative, attributable decision.

Two enforcement points, deliberately redundant. Visibility filtering at listing time means a principal never sees a tool outside its scope, so a model cannot be steered toward a tool it has never seen. The policy check at call time means a call outside scope or above the caller's ceiling is denied even if the caller somehow learned the tool's name. Neither point trusts the other. A defect in the first is contained by the second, and a caller probing names it should not know is denied and leaves a denial in the trail, which is itself signal.

3. Trust boundaries

There are exactly two trust boundaries, and the rule that separates them is data residency.

The customer boundary encloses everything that sees a raw tool-call payload: the gateway, the runtime that executes agents, every adapter wrapping a real system, the model path, and the stores holding full-fidelity audit events, run transcripts, and regulated data. All of it deploys into the customer's own cloud account as versioned artifacts, meaning container images plus an infrastructure construct, rather than source.

The Omnafy SaaS boundary encloses the multi-tenant governance core we operate: the policy plane holding principals and scopes and approvals, the control plane holding the definition registry and triggers and metering, the console, and a redacted audit store.

What crosses between them is small, enumerable, and directional.

DIRECTION	TRAFFIC	MUST NEVER CONTAIN
SaaS to customer	Approved definition versions and lifecycle state	Any other tenant's state
SaaS to customer	Policy state: principals, scopes, tier ceilings, graduation state	Any other tenant's state
SaaS to customer	Approval decisions, with approver identity	Any other tenant's state
Customer to SaaS	Redacted audit envelopes	Payloads, regulated data, transcripts
Customer to SaaS	Metering counters	Anything beyond counts and dimensions
Customer to SaaS	Health telemetry	Any field not affirmatively classified as safe

Everything else stays inside the customer boundary: tool-call payloads, adapter traffic to business systems, transcript capture, full-fidelity audit writes, and approval evidence.

One transport rule governs the seam and is easy to miss on a diagram. **Every boundary connection is initiated from the customer side.** Downward state is pulled by the customer stack from our replication endpoint. Upward events are pushed by the customer stack to our ingestion endpoint. We hold no credentials, no network path, and no cross-account role that reaches into the customer account. The single cross-account grant in the whole topology points the other way: a read-only pull on our container registry, exercised by the customer.

Two consequences follow that later sections depend on.

Policy enforcement is customer-side; policy authority is SaaS-side. The policy plane is authoritative for principals, scopes, tiers, and graduation state, and that state replicates down to the gateway, which evaluates every listing and every call locally. No per-call request crosses the boundary to get a decision, so no argument material does either. If the replica is unavailable or stale beyond its validity bound, the gateway fails closed.

Redaction happens at the boundary, on the customer side. The classification decision for an upward-flowing field is made inside the customer's account before transmission. We never receive material that we then have to promise to discard.

4. Identity

There is one identity model for every caller class. Human, agent, and service principals differ in how their identity is issued, not in how it is enforced. Nothing executes anonymously or under ambient authority: the gateway resolves every session to exactly one principal before any listing or call proceeds, and a session it cannot resolve is refused.

4.1 Human principals

Staff authenticate against the customer's own identity provider. We operate no staff identity store and never hold staff passwords. The principal record in our policy plane maps an identity-provider subject to a principal id, its scopes, and its tier ceiling. It holds no credential material.

A staff member's AI client, meaning any conforming MCP client, completes an OAuth 2.1 authorization-code flow with PKCE against the deployment's auth broker, which federates to the customer identity provider over OIDC and issues the gateway's own short-lived access token naming the human principal. The broker never handles identity-provider passwords, and the gateway validates each presented token locally against material in the replicated policy state, with no per-call round trip to the identity provider.

Three consequences are deliberate. The principal is the person, not the client and not the model: the same clinician using two different AI clients is one principal with one scope, one ceiling, and one audit history. Deprovisioning is inherited: disabling a person in the customer identity provider ends their ability to mint new sessions, and disabling their principal in the policy plane takes effect at the gateway on the next call, because every call is re-evaluated and there are no cached allow decisions. And what the client receives is governed while where the client sends it is sanctioned, which is the subject of section 7.3 and is the single most important thing for a security reviewer to understand about the entry product.

4.2 Agent principals

An agent principal is minted by the runtime at run start, bound to the pair of definition version and run id, and scoped to exactly what the approved definition requests. It is not a durable account. It cannot outlive its run, a new run of the same definition receives a fresh principal, and the kill switch halts the run and invalidates the identity with it. Because the identity carries the definition version and run id, every audit event an agent produces is attributable to the exact reviewed artifact that authorized the behavior, rather than to a generic automation account.

Sub-agents spawned within a run are not new principals. They execute under the parent run's own agent principal: same run id, same principal, same effective scope, same tier ceiling. Intra-run delegation therefore cannot escalate by construction rather than by check. An agent that wants different authority does not spawn. It drafts a definition, which crosses the full non-escalation machinery described in section 5.4.

The mint itself is bounded twice. The run's bound scope persists immutably on the run record, and the gateway resolves every call of that run at the intersection of the durable identity's current scope with that cap. Narrowing the durable identity narrows every live run, and the cap can never widen one.

4.3 Service principals

Service principals cover machine integrations that are neither people nor runs, such as the endpoint that receives inbound documents and fires triggers. They are issued by the policy plane, named per integration and never shared, granted the narrowest scope that lets the integration function, and rotated on a short schedule. A service principal that only enqueues trigger events holds no tool scope at all.

4.4 Token lifecycle

PRINCIPAL TYPE	ISSUED BY	BOUND TO	LIFETIME	REVOCATION
Human	The deployment's auth broker, federating to the customer identity provider	Person and session	Short-lived access token	Identity-provider deprovisioning stops new sessions; policy-plane disable denies on the next call

PRINCIPAL TYPE	ISSUED BY	BOUND TO	LIFETIME	REVOCATION
Agent	The runtime at run start, via the in-boundary token issuer	Definition version and run id	The run, capped by the run timeout	Kill switch; run completion or timeout
Service	Policy plane	One named integration	Short-lived, auto-rotated	Policy-plane disable, effective on the next call

Three rules apply across all three rows. There are no long-lived static credentials for any caller class, so anything bearer-shaped expires on its own. A token names exactly one principal and is never shared between people, runs, or integrations, because a shared token breaks attribution. And bearer material never enters the record: audit events and transcripts carry principal ids, never tokens.

4.5 Downstream credentials

The gateway and the adapters behind it hold credentials for real systems that are broader than any single caller's authority, because someone legitimately needs each operation. That is the classic confused-deputy setup, with the gateway as the deputy.

The defense is that downstream credentials are least-privilege and principal-scoped. An adapter executes a call using credential material scoped to the calling principal's authority, so a narrow caller can never borrow the gateway's full downstream reach by routing through it. A principal scoped to reads meets a read-only credential at the real system even if the adapter could, for some other principal, write. The gateway-side half of the same defense is that the principal context injected on dispatch comes from the gateway's own resolution and is never copied from anything the caller supplied.

4.6 Console and approver identity

The console is the highest-authority human surface in the platform. Approvers release held actions through it, graduations are signed off through it, and policy mutations originate from it. Its identity model is therefore stated rather than left implicit: console sign-in federates to the same customer identity provider as every other human access path, and a console session resolves to the same human principal record that a gateway session would. One person, one principal, one audit history, whichever surface they arrive through.

Multi-factor and session policy inherit from the customer identity provider, because the identity provider applies them at federation and we operate no weaker parallel path. Deprovisioning inherits identically. And the console holds no authority of its own: it is a client of the administrative API, so every mutation

and every release it performs is a tiered, audited operation attributed to the signed-in principal, worth exactly that principal's scope and ceiling. A hijacked console session is analyzed as stolen-credential misuse, not as a separate super-user surface.

5. Authorization

Authorization answers three questions for every call, in order. May this principal see this tool at all, which is scope. May it invoke a tool of this risk class, which is the tier ceiling. May it act without a human in the loop right now, which is the autonomy posture.

5.1 Scopes

A scope is a named grant with three components: selectors over namespaced tools and resources, a tier ceiling, and tenant slices. A principal's effective scope is the resolution of all its grants into a single comparable value, and that value is what both enforcement points consult.

The grammar is deliberately impoverished. A tool selector is either an exact namespaced name or a prefix with a single trailing wildcard. There are no mid-pattern wildcards, no negation, no regular expressions, and no conditional predicates. There are no negative grants at all: a scope only ever adds authority, which keeps composition monotone, so the union of two grants is a well-defined grant and adding one can never silently revoke what another gave.

That restriction is not aesthetic. It is what makes subset comparison exact, which section 5.4 explains.

5.2 Risk tiers

A risk tier is a property of the tool, not of the caller. A tool's tier answers what the worst plausible consequence of this operation is, which is invariant across callers. The caller contributes only a ceiling.

The tier is declared in the tool's manifest and ratified by a human at registration. We never take a tool server's word for its own risk. The four tiers form a total order.

TIER	NAME	COVERS	DEFAULT AUTONOMY POSTURE
T0	Observe	Read-only monitoring, reporting, retrieval	Autonomous immediately
T1	Internal actions	Reversible writes inside customer systems	Propose then approve, until graduation is earned per principal and task
T2	External actions	Anything leaving the boundary: external records, faxes, messages	Approval-first for longer; graduates selectively
T3	Judgment calls	Professional or legal judgment	Permanently human. Agents recommend with evidence; T3 never graduates

Tier and posture are orthogonal checks. Being within ceiling admits the call class; the autonomy posture for the pair of principal and task then decides between allow and pending-approval. Graduation changes only the posture, never the ceiling.

T3 deserves a paragraph of its own because it is the platform's most load-bearing promise. T3 is not a slower T2. Ratifying a tool at T3 is the organization stating that no history of agreement, however long, converts this decision into an automation target. Its permanence is enforced in three independent places. The policy plane has no representable graduation state for a T3 tool, so there is nothing to misconfigure. Non-human principals are ceiling-capped at T2 by construction, so a corrupted policy state that invented a T3 graduation entry would still meet a caller that cannot clear the ceiling check. And visibility filtering never lists a T3 deciding form to a non-human principal, so a model cannot be steered toward a tool it has never seen.

5.3 The two enforcement points

Visibility filtering at listing time. The gateway computes each principal's visible tool list from its effective scope before the response is built and before any model sees it. A catalog entry survives into a listing only if it is pinned and ratified, servable, within the principal's scope selectors, under the principal's tier ceiling, compatible with the principal's tenant slices, and, for a T3 deciding form, being listed to an eligible human. Filtering strictly precedes pagination, so no cursor boundary can leak the existence of an out-of-scope tool.

The policy check at call time. Every call is independently evaluated against the replicated policy state: scope membership, tier ceiling, tenant compatibility, manifest integrity, an active kill switch covering the calling run, a declared task for agent callers, and the autonomy posture that applies. Any input the gateway cannot evaluate produces a fail-closed denial.

Calling an invisible tool by name returns an error byte-identical to calling a nonexistent one, while the attempt is fully audited as a denial. Invisible to the caller, visible in the trail.

5.4 Non-escalation

Invariant I3 is the answer to the question every buyer eventually asks: what happens when an agent writes an agent?

An agent may draft a definition. Drafting is a propose-class operation like any other. What it may not do is draft a definition whose authority exceeds its own. The check is a mechanical subset comparison over effective scopes:

$S_1 \subseteq S_2$ holds when $\text{selector-set}(S_1) \subseteq \text{selector-set}(S_2)$ and $\text{ceiling}(S_1) \leq \text{ceiling}(S_2)$ and $\text{slices}(S_1) \subseteq \text{slices}(S_2)$.

Because selectors are exact names and trailing-wildcard prefixes, selector inclusion reduces to string and prefix containment over canonical forms. It is decidable, cheap, and exact.

Three properties make the check hold up against the interesting attack, which is escalation laundered through composition. First, the comparison runs over the closure of everything a composition can reach, not just the draft's directly requested scope. An agent holding intake scope that drafts a workflow whose steps invoke a chart-assembly agent with broader clinical scope is caught, because the chart-assembly agent's scope is in the operand. Second, the check runs twice: at draft time, and again at approval against the creator's effective scope as of approval, so a grant that narrowed in between cannot be smuggled past on a stale pass. Third, two scopes that cannot be compared are treated as failing the check and the draft is rejected. Incomparability fails closed.

Independently of all three, only a human may approve a definition. An agent can draft; it can never activate.

The grammar-extension rule that protects this is stated as a standing decision: any future extension may only further restrict an allowance after the subset check, so the unconditional scope remains a sound over-approximation of authority. An extension that makes the comparison undecidable or heuristic does not ship.

5.5 Fail closed, concretely

Each of the following produces a denial rather than a default allowance:

The session cannot be resolved to a known, enabled principal. Replicated policy state is unavailable or stale beyond its validity bound. The tool has no manifest entry, or its manifest hash no longer matches what was ratified at registration. The manifest carries no ratified tier. The scope grammar cannot evaluate the comparison, for example against state produced by a newer grammar version than the evaluator understands. An approval queue item expires without a decision, in which case expiry escalates or rejects and never auto-approves.

Availability is deliberately subordinate to this rule. A gateway that cannot reach policy state stops serving tools rather than serving them ungoverned. Section 8.4 states the cost of that choice plainly.

6. The record

Invariant I4 says every action carries the full principal chain into an append-only audit event. This section describes what that record actually is, because it is the artifact an auditor will query and the reason the other four invariants are checkable rather than merely claimed.

6.1 The envelope

Every event carries a core envelope: a schema version, a globally unique event id that doubles as the ingestion idempotency key, an event type, a timestamp, a tenant id, the emitter stream, a strictly monotonic per-stream sequence number, the SHA-256 hash of the previous event on that stream, an optional correlation id, the principal chain, and, per type, blocks for the tool, the policy decision, the run, the payload, the outcome, and the cost.

The per-stream hash chain and the monotonic sequence are the tamper-evidence mechanism. Sequence gaps alarm, and modification of a written event breaks the chain.

6.2 The principal chain

The chain is the envelope's attribution spine and is mandatory on every event. It has three roles in order.

`caller` is the principal that performed the action. For agent callers it also names the run, so attribution reaches the exact execution rather than the durable identity.

`onBehalfOf` is the principal whose delegated authority the caller is exercising. A run carries the human principal whose act put its definition live or who invoked it manually. A trigger-fired run carries the principal that approved the trigger binding. An interactive staff call has no delegation and omits the field.

`approver` is the human principal whose approval released this action, with the queue item id, present exactly when the action executed under a release.

The chain is what makes the trail usable for confused-deputy forensics. Because every event records the narrow caller whose authority was exercised, an investigator can join any effect observed in a business system back to a specific chain and ask one mechanical question: was this effect within the caller's effective scope at the recorded policy state version? An effect with no matching event is a suppression finding. An event whose effect exceeds its caller's recorded authority is the forensic signature of a deputy acting on its own credentials. Without the chain, the trail could only say that the adapter did it, which is true of every action.

Denials carry the same chain. The record has no allow-only bias. A caller probing tool names it should not know leaves a sequence of attributed denial events, which is the detection surface for stolen-credential misuse and for injection-driven probing.

6.3 The seven event types

TYPE	EMITTED BY	RECORDS
tool-call	Gateway	One dispatched tool call and its result, including administrative operations invoked as tools
policy-decision	Gateway, both enforcement points	Every denial including every fail-closed denial, every pending-approval hold, and listing-time visibility filtering
approval-decision	Policy plane	The resolution of one queue item: released, rejected, or expired
lifecycle	Control plane	A definition lifecycle transition, including the recorded result of the non-escalation check
graduation-change	Policy plane	A promotion, demotion, certification pass, or lapse for one principal and task pair
admin-mutation	Control plane	Any other governed mutation of policy or configuration, including a kill-switch exercise
registration	Gateway and control plane	A tool server joining or changing: registration, namespace assignment, tier ratification, collision rejection, manifest change

Every verdict from either enforcement point is recorded exactly once. An immediate allow travels inside the `tool-call` event it authorizes, so an executed call maps to exactly one audit event. A deny, a hold, and a filtered listing have no execution to ride on and stand alone. Nothing is recorded twice and nothing is recorded zero times.

6.4 Two materializations, one event

Every event exists in up to two forms. The **full-fidelity record** lives in the customer boundary and carries argument and result material. The **redacted envelope** crosses to our control plane and is byte-identical except that the payload block retains only keyed hashes and the outcome block drops its free-text detail slot.

The shared event id and the keyed payload hashes bind the two. A customer holding the HMAC key and the full-fidelity record can recompute the hashes and confirm that the envelope we hold corresponds exactly to the event they hold, with no payload ever crossing the boundary.

The hashes are keyed for a specific reason. A plain digest of a low-entropy field, say a date of birth or a phone number, is invertible by dictionary. An HMAC under a customer-held key that never leaves the customer account is not, and it also cannot be correlated across tenants.

6.5 Recording gates execution

One guarantee is load-bearing and is deliberate coupling: **an action the gateway cannot record does not happen**. The full-fidelity audit write is durable before the result returns. Customer-side audit store unavailability therefore suspends tool execution rather than degrading the trail. Fail closed, applied to recording.

The upward channel is treated differently on purpose. If it is down, calls are unaffected: envelopes buffer with sequence numbers intact and redeliver. The customer-side trail is canonical and gates work; the upward copy is a governed export and never does.

6.6 Storage properties

The stores are append-only, using object-lock compliance mode on both sides. The write path is separated from readers and operators, so no role that can query the trail can rewrite it. Integrity checking detects modification of what was written. Corrections are new events carrying a `supersedes` reference, never edits.

Suppression is countered by the dual-record design. Every action emits both a customer-side full-fidelity event and a SaaS-side redacted envelope, so suppressing either stream creates a detectable mismatch with the other, and alarms fire on gaps and sequence discontinuities rather than waiting for a human to notice an absence.

7. Data residency

Every category of data the platform touches has an assigned residency: where it may live and, the operative property, where it must never appear.

DATA CLASS	MAY LIVE IN	MUST NEVER APPEAR IN
Regulated data	Customer business systems; adapter traffic; transcripts; full-fidelity audit; the customer's model path	The Omnafy control plane in any form, including the redacted audit store, metering, telemetry, and logs
Tool-call payloads and transcripts	Customer-side stores only	Anything crossing the deployment boundary
Tool results delivered to staff clients	The staff member's AI client and the model path the customer has sanctioned behind it	The Omnafy control plane
Full-fidelity audit events	Customer-side append-only store	The Omnafy SaaS
Redacted audit envelopes	Generated customer-side, stored in the SaaS	Constrained by content, not location: no payload fields, no regulated data

DATA CLASS	MAY LIVE IN	MUST NEVER APPEAR IN
Policy state	Authoritative in the SaaS, replicated to the gateway	Contains no regulated data by construction
Definitions	SaaS registry, replicated to the customer runtime	Must never embed regulated data or credentials, enforced at draft validation and again at approval
Secrets and credentials	Customer-side secret store; SaaS-side store for control-plane service credentials	Definitions, transcripts, audit events, telemetry, source control
Metering and health telemetry	Crosses upward	Any field not affirmatively classified as safe

7.1 The redaction pipeline

One component in the customer stack holds credentials to publish to our ingestion endpoint, and every upward byte passes through it. It enforces a deny-by-default field policy with three rules.

Classification is affirmative and versioned. A field crosses only if it appears in the field classification table with a safe verdict. The table is a reviewed artifact shipped with the customer stack. A field absent from the table is dropped, and the drop is counted in health telemetry, so a new unclassified field surfaces as a signal rather than as a silent leak or a silent loss.

Free text never crosses. Error messages are mapped to enumerated error classes before transmission, because a raw vendor error string can embed regulated data. Identifiers cross; prose does not.

Payload hashes are keyed, as section 6.4 describes.

The failure mode this inverts is the usual one. Under deny-by-default, an unclassified field breaks telemetry rather than confidentiality. The remaining exposure is a field affirmatively misclassified as safe, and field-classification changes are reviewed changes on the audited configuration surface.

7.2 The control plane is a low-value target for data theft

A full compromise of our SaaS yields tenant metadata, policy state, definitions, and redacted envelopes. No regulated data, no payloads, no transcripts. That is a deliberate outcome of I5 rather than a side effect.

It remains a high-value target for authority, which is why policy mutations are themselves tiered, audited operations crossing the same gateway as everything else, and why cross-tenant leakage is treated as a first-class threat in section 8.

7.3 The one flow the platform cannot constrain

The largest regulated-data egress channel in the entry product is the governed one, and we state it rather than eliding it.

When a scoped staff principal reads records through a desktop AI client, the result travels to that client and whatever model backend sits behind it. It leaves the customer's cloud account lawfully, by design, because delivering the result is the point of the product. Invariant I5 constrains our infrastructure and the automation model path. It does not, and cannot, constrain what happens inside a client the customer has sanctioned to receive results.

The platform bounds this flow with two mechanisms. Scope decides what a principal can obtain at all, which is a minimum-necessary control. And every delivery is on the record, so what was delivered to whom is queryable. The customer bounds it with a third: sanctioning only clients whose model path meets its compliance bar. That is a named customer responsibility in the shared-responsibility matrix, not fine print, and the corresponding agreement sits in the customer's contract chain rather than ours.

A security reviewer should treat this as the first question to ask about the entry product, and we would rather they read it here than discover it in month three.

8. Threat model summary

The analysis is a STRIDE pass over the trust-boundary and container diagrams. For each node and each edge that crosses or touches a boundary, we asked which categories an adversary could realize there and kept every result that survives contact with the invariants. The output is not an exhaustive matrix. It is the ten threats that matter, ranked by how directly they attack what the platform promises.

The document has one contract with its reader: every threat resolves to either a named, documented mechanism or an entry in the accepted-risks list. There is no third category. A threat that resolved to "we are careful" would be a defect.

Six adversary classes drove the enumeration: an external network attacker, the author of adversarial content entering through legitimate channels, a malicious or compromised tool-server author, an attacker holding stolen principal credentials, a co-resident tenant, and an Omnafy insider. The last is handled under accepted risks.

8.1 The ten threats

#	THREAT	CATEGORY	MITIGATING MECHANISM	RESIDUAL
1	Prompt injection via tool results	Elevation of privilege	Tier gates, approval-first defaults, structural T3 permanence, listing-time visibility filtering. Never prompt hygiene alone	Monitored
2	Confused deputy at the gateway	Elevation of privilege	Per-principal scoping of downstream credentials plus an independent call-time policy check	Low
3	Agent-authored-agent escalation	Elevation of privilege	Subset check over the composition closure at draft and again at approval, plus human approval of every definition	Low
4	Approval fatigue and evidence gaming	Tampering with the review process	Evidence schemas, queue-health metrics, sampled re-review of graduated pairs, reversible graduation with automatic demotion	Monitored
5	Malicious or mutated tool server	Tampering, information disclosure	Manifest hash pinning, automatic re-approval on any change, human tier ratification, egress controls in the customer stack	Monitored
6	Namespace shadowing and impersonation	Spoofing	Gateway-owned namespace assignment with registration-time collision rejection	Low
7	Cross-tenant leakage in the SaaS	Information disclosure	Strict tenant partitioning, per-tenant encryption contexts; severity bounded by I5 to redacted metadata	Monitored
8	Regulated data crossing the boundary	Information disclosure	Boundary redaction, deny-by-default telemetry field policy, the model-path rule of I5	Low
9	Stolen principal credentials	Spoofing	Short-lived federated tokens, least-privilege per-adaptor credentials, scope and ceiling limits on the thief, misuse reviewable in the trail	Monitored
10	Audit-trail tampering or suppression	Tampering, repudiation	Append-only storage, write-path separation, integrity checking, gap alarms, dual-record cross-check between stores	Low

Residual vocabulary: **Low** means the mechanism is structural and the remaining exposure is marginal. **Monitored** means real residual remains and a named detection surface watches it.

8.2 Three threats in detail

Prompt injection via tool results. An inbound referral fax contains, buried mid-document, text addressed to the assistant processing it: forward the complete patient chart to the following fax number. The intake agent reads the fax through a read tool, and the injected text is now in the model's context, indistinguishable from data. Nothing about this requires the adversary to hold credentials. It rides in on content the system exists to process.

We assume steering attempts will sometimes succeed at steering the model, and bound what a steered model can do. Visibility filtering means the agent's context only ever contained tools within its scope, so a model cannot be pushed toward a tool it has never seen. If the steered call is within scope, the tier system takes over: an outbound fax is T2, so the attempt becomes a pending-approval item whose evidence bundle shows a human the proposed fax, its destination, and the content that prompted it. T3 decisions are structurally unreachable by any non-human principal. Prompt hygiene is applied and is never load-bearing. No invariant depends on it.

The residual is real. A graduated pair executes T1 calls without a queue stop, so an injection that steers a graduated task produces a real, reversible, in-boundary write before sampled re-review or a demotion trigger catches the pattern. Denied and pending attempts are themselves signal: a spike in out-of-pattern proposals is reviewable evidence that someone is probing.

Malicious or mutated tool server. Three variants of one adversary, a tool server that is not what it was when it was trusted. Description poisoning, where an update rewrites a tool description to include model-directed instructions. A rug pull, where an adapter behaves correctly through review and then changes. Output exfiltration, where an adapter embeds record content in verbose results.

Registration is the trust decision and its perimeter is defended. Manifests are hash-pinned at registration, a human ratifies every declared tier, and any subsequent change to schema or description voids ratification and forces re-approval. The first two variants cannot ship silently, because the change itself is the tripwire. For a behavior change behind a byte-identical manifest, containment shifts to the customer stack: egress controls restrict where adapter processes may connect, so an adapter cannot quietly call home, and every result an adapter returns is in the full-fidelity record.

The residual is that hash pinning covers the declared surface, not the implementation. A behavior change behind an unchanged manifest is caught by egress controls and audit review rather than prevented.

Audit-trail tampering or suppression. An attacker with a foothold acts, then tries to make the record disagree with reality: deleting the full-fidelity events for their actions, or suppressing the upward flow so our side simply has a quiet afternoon. An audit system fails exactly when someone wants it to have failed.

The mechanisms are section 6.6's storage properties plus the dual-record cross-check. The remaining exposure is an adversary who fully controls the customer-side emission point at the moment of action, at which point they are inside the gateway itself, which is the perimeter every other control also assumes holds. Cross-checking against the independently stored envelopes still bounds what such an adversary can silently rewrite after the fact.

8.3 One named data flow that no mechanism makes safe

A curated catalog may include an entry served by a vendor-operated remote endpoint outside both trust boundaries. That endpoint receives full tool-call payloads.

The flow is declared on the catalog entry, registration requires the tenant administrator's explicit recorded acceptance, egress to the vendor host widens only as a consequence of ratified registration state, vendor authentication material is a registered secret in the customer-side store, and pinned manifests are drift-checked against the endpoint.

What none of that changes: the vendor sees the payloads. The catalog makes the flow declared, ratified, and audited. It cannot make it safe. The vendor agreement is the customer's to hold. We are not party to it.

8.4 Accepted risks

These are the exposures we currently accept rather than mitigate. The list is short and explicit. An exposure that belongs here and is not here is a defect to raise.

No per-tenant global rate governor. Per-run budgets, timeouts, and spawn depth limits bound the cost of any single runaway run, and the kill switch halts one on demand. There is not yet a per-tenant global governor, so a trigger storm could start many individually-within-budget runs whose aggregate cost is painful. Accepted because the exposure is financial and operational rather than a confidentiality or integrity break, per-run bounds cap the slope, and metering makes the aggregate visible quickly. Closed by a per-tenant rate and spend governor.

The scope algebra is tested, not formally verified. Invariant I3 rests on the subset comparison. The grammar was deliberately restricted to make that comparison exact, and the canonicalization and comparison code is tested, but there is no machine-checked proof that canonicalization preserves the semantics the comparison assumes. Accepted because the grammar's restrictions keep the algebra small enough that testing covers it densely, and the human-approval gate on every definition means a failure requires a human miss as well. Closed by property-based exhaustive testing over the grammar, or formal verification of the canonicalizer and comparator.

Insider risk at Omnafy, pending formal controls. Our operators administer the SaaS control plane and could, absent further controls, mutate policy state or definitions. Invariant I5 bounds what an insider can take, since the control plane holds redacted metadata and no regulated data, but authority over policy state is itself sensitive. Accepted because the team is presently small enough for organizational mitigation, and because administrative mutations are themselves tiered, gateway-crossing, audited operations that land in the same trail customers can query. Closed by formal change-management, access-control, and monitoring controls on the SOC 2 readiness path.

Gateway availability. The gateway is deliberately a single mandatory path, fail-closed everywhere. Every executed call couples to the customer-side audit store, policy-replica staleness denies, and a gateway outage stops all governed AI access and all automation at once. No formal availability objective, redundancy attestation, or continuity test exists yet. Accepted because the failure mode is the safe one by construction: nothing fails open, no ungoverned fallback path exists to inherit the traffic, and the customer's manual process remains available, since the platform governs AI access to systems of record and is not the systems of record. Closed by availability objectives, redundancy, and continuity testing for both stacks.

Publishing this list costs a little polish. It buys the thing this reader actually purchases, which is the confidence that our claims and our architecture are the same document.

9. Implementation status

Design documents describe an intended system. This section says which parts of the intended system exist, because a security review that cannot tell the difference is not a security review.

Implemented. The governance core is built: the gateway with its aggregation, namespacing, manifest pinning, visibility filtering, and call pipeline; the policy plane with scopes, tiers, the subset comparison, approval queues, and graduation state; the control plane with the registry, definition lifecycle, triggers, the dual-exposed administrative API, tenancy, and metering; the runtime with run lifecycle, transcript capture, budgets, spawn depth limits, and the kill switch; the redaction pipeline; and the audit pipeline across both stores. The automation module has an end-to-end path under test that drives a scheduled firing across the boundary channel into a real run whose tool calls reach a real adapter through the real gateway, parks it on a hold, releases it through an approver's decision, and asserts the record set on both sides of the boundary.

Designed, not shipped. Communications adapters for email, SMS, and fax are designed at T2 by default, because their actions leave the boundary. Until they ship, outbound communication steps in any workflow are human-executed. A browser-automation adapter tier for systems with no API is designed at T2 with an additional verification layer required at registration, and that is the only form in which browser automation will ever exist in the product. Cross-tenant isolation tests in continuous integration are pending, which is why threat 7 carries a Monitored residual rather than a Low one. Automated caller-behavior anomaly detection is pending, which is why detection latency for threat 9 is currently the latency of the customer's review practice.

Not claimed. No SOC 2 report exists and no auditor has been engaged. No independent penetration test has been performed and no assessment report exists to hand over. There is no formal availability objective. Section 8.4 states each of these as an accepted risk rather than leaving the reader to infer it.

10. What this architecture does not claim

It does not claim that models behave predictably. It claims that a model's behavior is bounded by scope, tier, and posture, and that the bound holds whether or not the model cooperates.

It does not claim the absence of residual risk. Section 8.4 is the list, and it changes only by shipping a mechanism, never by optimism.

It does not claim that supervision cannot decay. It claims that decay is visible in queue-health metrics, that it is never final because sampling and demotion operate after supervision ends, and that the customer who ignores their own dashboard is the failure mode that remains.

It does not claim to constrain the model path behind a client the customer sanctions. Section 7.3 states what it does instead.

Further reading

The companion papers in this series go deeper on the surfaces this one summarizes. *The Omnafy platform architecture* covers the five planes, the gateway internals, the contracts, and the split deployment in engineering detail. *Earned autonomy* covers risk tiers, approval queues, and the graduation mechanism from the operations side. *HIPAA and the AI governance layer* covers the shared-responsibility split, the agreement chain, retention, deletion, and our SOC 2 readiness gap.

Customers and prospective customers can request the full internal documentation set, including the complete threat model with all ten scenarios written out, the audit events contract with its schemas, and the decision log.

Glossary

Adapter. A tool server whose purpose is to wrap a specific real system and expose its operations as governed tools.

Definition. A declarative, versioned artifact describing automated behavior, either an agent definition or an explicit workflow. Lifecycle: draft, approved, live, retired.

Effective scope. The resolved authority of a principal at a moment: all its grants evaluated to a single comparable value.

Evidence. The material attached to a proposed action that lets a human decide it on the record: inputs, retrieved context, the proposed action, expected effects, and cost.

Gateway. The single aggregate MCP endpoint through which all tool traffic flows.

MCP. Model Context Protocol, the open protocol connecting AI clients to tool servers. We did not invent it. We govern it.

Manifest. The per-tool declaration a tool server submits at registration: schema, risk tier, side-effect class, reversibility, evidence hints, cost hints, idempotency, dry-run support.

Namespace. The gateway-owned prefix that partitions tool names, assigned at registration and never chosen by a tool server.

Principal. The identity every caller carries. Three types, one model: human, agent, service.

Principal chain. The ordered attribution on every audit event: who acted, on behalf of whom, approved by whom.

Redaction. The boundary pipeline that strips regulated data from anything crossing up into our control plane, under a deny-by-default field policy.

Risk tier. The risk classification of a tool, declared in its manifest and ratified by a human at registration, never asserted by the caller.

Run. One execution of a definition, with its own run id, agent principal, transcript, budgets, and metered cost.

Scope. A named grant with selectors, a tier ceiling, and tenant slices.

Tier ceiling. The maximum risk tier a principal may invoke, carried as a component of each of its scopes.

Tool server. A service implementing the tool-server contract. Every adapter is a tool server; the administrative API is a tool server that is not an adapter.

Transcript. The complete record of a run: every message, model interaction, and full payload. Stays in the customer's account.