

AI Scaffolding

An Enterprise Control Plane for Safe, Scalable Adoption

By

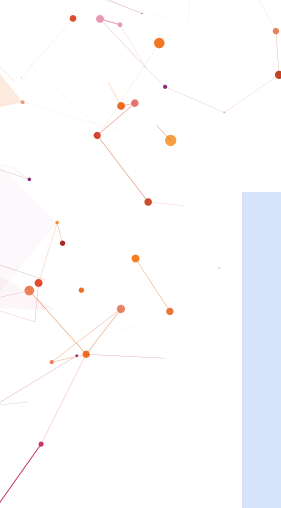
Mark Baker

Senior Growth Partner, Altimetrik

Frankie LaCarrubba

Architect - Product and Platform Engineering





I. Executive Summary

AI adoption is advancing faster than any enterprise technology in history. The companies pulling ahead are those boldly empowering their businesses to adopt and innovate at full speed. That boldness is delivering results but also raising the stakes. The broader business community is waiting for that seminal moment: the AI equivalent of an evening news story about a lost laptop exposing gigabytes of customer data – proof that too much freedom can lead to disaster.

The way forward is a new architectural construct we call **AI Scaffolding**. AI Scaffolding extends the traditional enterprise control plane with the services, guardrails, and automation needed to sustain rapid adoption, turning governance into the mechanism that ensures speed, safety, and scale.

What AI Scaffolding delivers:

- **Visibility** - Registries for GPTs, agents, tools, and prompts, capturing ownership, purpose, and evaluation status, and securing AI assets as first-class governed assets.
- **Governance** - Risk-tiering and policy frameworks aligned with NIST AI RMF, ISO/IEC 42001, and (*wherever applicable*) the EU AI Act, enabling proportionate controls that keep low-risk adoption fast and lightweight.
- **Enforcement** - An Agent & MCP Gateway plus Guardian Agents that enable automatic policy enforcement in real time.
- **Interoperability** - Standards (*Agent Cards, MCP connectors*) that prevent fragmentation and vendor lock-in.
- **Efficiency** - shared infrastructure, observability, and FinOps controls delivered once, centrally.

With Scaffolding in place, enterprises can move faster than competitors, not by ignoring risk but by mastering it, and thus defining what true leadership in enterprise AI adoption looks like.



II. Context and Rationale

AI adoption today is no longer optional. It is being driven both by visionary leadership that recognizes its potential and by competitive pressures that leave no alternative. Some of our customers have already leaned further than their peers, advancing at a pace more typical of high tech than traditional industries. That boldness is propelling them ahead — in both adoption and value creation.

Yet history reminds that speed without structure brings its own risks. Red tape rarely exists without reason. Across industries, many are waiting for “the other shoe to drop” — a high-profile incident that proves too much freedom can be reckless. We see this as a moment to prove the opposite: that with the right approach, governance not only

keeps organizations out of the headlines but also sustains adoption velocity and amplifies innovation.

Governance teams face a parallel mandate just as urgent as adoption itself: to manage AI's risks in ways that don't just permit rapid use but actively enable it. That requires a lighter touch like controls that accomplish more with little or no human intervention, and a unified approach where policy and enforcement are simultaneous. The goal is to keep the floodgates wide open when they can be, and snap them shut instantly when they must be.

Enterprises have long relied on governance and control planes to keep technology adoption safe and consistent. But AI changes the equation. Traditional approaches

cannot match the velocity and scale this new wave demands. Governance must now operate with sharper discernment and unprecedented speed. This is where the concept of Scaffolding comes in.

In the physical world, Scaffolding is lightweight, rapidly deployed, and easy to adapt, while giving workers safe access to an entire structure. We use the term here to describe an architectural layer that extends the traditional enterprise control plane with additional services, guardrails, and automation that act quickly, adapt in real time, and give the business access to the full leverage of AI tools at maximum speed.

With Scaffolding in place, adoption can move as fast as ambition demands, transforming governance from a brake on innovation into

III. Decentralized AI Challenges

Rapid AI adoption brings enormous leverage. When each business unit moves to the beat of its own drummer, that leverage multiplies risks instead of value. The benefits of freedom quickly turn into fragmentation. Without Scaffolding, enterprises racing ahead face several predictable pitfalls:

- **Shadow AI usage** – Agents, tools, and prompts proliferate without shared standards or enterprise-wide visibility.
- **Inconsistent safeguards** – Sensitive data and risks are handled differently across teams, with some applying strong protections while others apply none.
- **Duplicated effort** – Logging, monitoring, and data controls are repeatedly rebuilt within individual projects instead of being delivered once as shared services.
- **Brittle prototypes** – Siloed proofs of concept that can't scale beyond a single team, collapsing under enterprise-level demands.

AI Scaffolding mitigates these challenges by providing centralized services and guardrails, with maximum autonomy. This allows decentralized teams to bring new assets to market quickly, without duplication or unmanaged exposure.



IV. Platform Strategy & Architecture Vision

With these challenges in view, AI Scaffolding takes shape as an AI control plane—a shared layer of services, standards, and guardrails that bring both agility and assurance to enterprise AI adoption.

The key elements are as follows:

Foundational Scaffolding

- **Identity and Access**

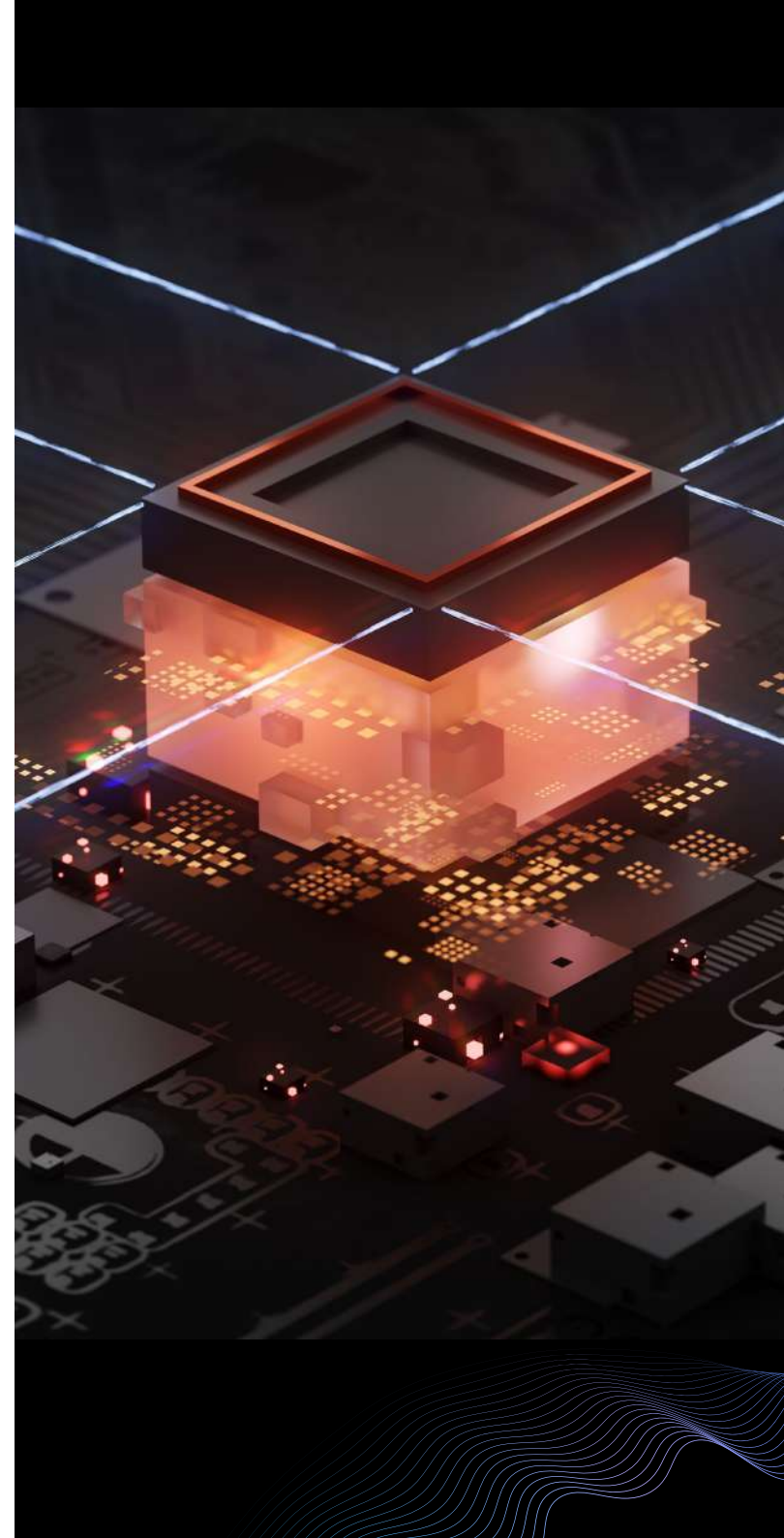
Service principals, RBAC/ABAC models, and per-agent secrets ensure that every GPT, agent, and tool is authenticated, authorized, and governed.

- **Registries for Visibility**

Centralized catalogs of GPTs, agents, MCP tools, and prompts — treated as governed assets with clear metadata such as owner, purpose, data access, and evaluation status — enable transparency and lifecycle management.

- **Governance Frameworks**

Risk-tiering, approvals, logging, and incident response processes align with standards such as **NIST AI RMF** and **ISO/IEC 42001**, building trust across compliance, legal, and security stakeholders.





Core Enforcement *(Runtime and Integration)*

• Agent & MCP Gateway

A central enforcement and routing layer between AI agents/tools and enterprise systems that:

- ◆ Enforce policies at runtime *(moderation, rate limiting, quotas, logging)*.
- ◆ Injects required metadata *(user, purpose, context)* for traceability.
- ◆ Standardizes ingress/egress for tools and agents via MCP servers and Agent Cards.
- ◆ Provides a single point of observability *(OpenTelemetry traces, SIEM/APM integration)* and secures access to sensitive data *(PII, PCI, PHI, confidential business data)*

• Agent & Tool Interoperability

Standardized connectors (**MCP**) and metadata (**Agent Cards**) enable safe, auditable interaction across systems while avoiding vendor lock-in.

• Region-Specific Architecture

Region-aware deployment of agents and MCP servers ensures compliance with data residency and sovereignty requirements.

Advanced Enforcement *(Sensitive or Edge Cases)*

• Walled Gardens for Sensitive Workloads

Isolated environments for high-risk domains such as audit, HR, and finance, feature enhanced logging, restricted data access, and additional monitoring controls.

• Multi-Model and Provider Abstraction

A pluggable architecture supports multiple LLM providers, allowing prompts, agents, and workflows to run seamlessly across models without re-engineering.

• Outcomes and Strategic Benefits

With these architectural elements in place, enterprises can:

- ◆ Encourage rapid experimentation without sacrificing visibility or governance.
- ◆ Equip builders with safe templates and tools, accelerating delivery of new use cases.
- ◆ Give governance teams the ability to limit exposure proactively or prevent it altogether.
- ◆ Align tightly with evolving regulatory requirements without slowing innovation.



V. Early Priorities

To establish AI Scaffolding effectively and accelerate adoption, enterprises should focus on a set of foundational priorities that balance visibility, governance, and speed of deployment.

Agent & MCP Gateway

Stand up a central enforcement and routing layer that:

- ◆ Secures sensitive data (*PII, PCI, PHI, and confidential business information*).
- ◆ Enforces runtime policies such as moderation, rate limiting, quotas, and logging.
- ◆ Provides unified observability through integrations with OpenTelemetry, SIEM, and APM systems.

This gateway serves as the central control point for real-time policy enforcement and traceability across all AI interactions.

Risk-Proportionate Governance

Differentiate high-risk from low-risk AI use cases. Apply proportionate controls — approvals, human-in-the-loop review, and detailed logging — only where necessary, keeping lightweight adoption fast and frictionless.

Tool Governance via MCP Connectors

Standardize how agents connect to enterprise systems through the Managed Connector Protocol (MCP):

- ◆ Define least-privilege scopes (*read vs. write, dataset boundaries*).
- ◆ Establish approval workflows for new connectors.
- ◆ Capture metadata such as tool owner, purpose, and risk tier in registries.
- ◆ Enforce runtime controls through the gateway, ensuring only approved and registered tools are callable.

This approach enables safe interoperability with enterprise platforms (*e.g., Salesforce, ServiceNow, HR, finance*) while preventing shadow integrations.

Discovery of Unsanctioned AI Use

Incorporate mechanisms to detect, flag, and catalog AI activity occurring outside sanctioned environments. Early visibility allows governance teams to close gaps before they evolve into systemic risk.

Standards Alignment

Ensure Scaffolding practices directly align with recognized frameworks such as the NIST AI RMF, ISO/IEC 42001, and the EU AI Act (where applicable), building external credibility alongside internal control.



VI. Architecture Framework Discussion

Scaffolding should align with both who consumes AI and how they access it. A coherent framework ensures adoption moves at speed, but within trusted pathways, with governance embedded at every step.

Three-Layer Consumption Model

A three-layer model ensures every group, from builders to business users to autonomous agents, to innovate quickly while staying within clearly defined boundaries.

Builders

(Advanced Users & Developers)

Builders create new agents, connectors, and prototypes.

AI Scaffolding provides:

- ♦ **Safe defaults and reusable templates** for faster development.
- ♦ **Evaluation pipelines and traceability** for transparent governance.
- ♦ **Centralized gateways** that handle enterprise-grade requirements automatically, enforcing PII redaction, logging, rate limiting, and sensitive data protection, so builders don't have to code these controls from scratch.

Guardrails ensure that what they build can scale securely to production without unmanaged risk.



Business Users

(Everyday Employees)

Business users generate assets and outcomes through approved, governed tools — for example, producing customer-facing content, creating custom GPTs, or automating workflows.

Scaffolding supports them with:

- ♦ **Intuitive, governed portals** that embed privacy, security, identity, and data-sensitivity controls.
- ♦ The ability to focus on outputs, without navigating technical complexity or worrying about compliance.

This allows everyday employees to safely harness AI capabilities while maintaining trust and accountability.

Extended Workforce

(Contractors & Service Providers)

Non-FTE users operate with **tightly scoped, least-privilege access**, yet can still perform agentic tasks safely.

Scaffolding enforces **runtime limits, logging, and policy checks** to prevent overreach while enabling contribution.

Examples:

- ♦ A contractor using an AI agent to draft a project plan, gated so only employees can approve and publish.
- ♦ A vendor support specialist using an agent to triage service tickets — confined to a safe knowledge base and barred from accessing customer PII.

Access Models

To support all user groups, Scaffolding provides two primary access pathways:

- ♦ Conversational interfaces (LLM/chat) for everyday users.
- ♦ Developer portals for builders, offering registries, orchestration tools, templates, and secure sandboxes.

By differentiating entry points, governance applies consistently at the “front door” — whether a user is experimenting with a chat agent or developing production-ready workflows.

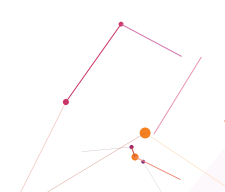
Embedded Controls

Supporting diverse roles requires deeper architectural capabilities:

- ♦ **Memory & Data Management** — Tiered memory (short-, mid-, and long-term), retention and compaction policies, and registry-inherited versioning for full auditability.
- ♦ **Guardrails Across the Lifecycle** — Input validation, schema enforcement, jailbreak and toxicity detection, hallucination checks, malware scanning, and DLP enforcement.
- ♦ **Orchestration & Multi-Agent Workflows** — Coordinated task handoffs, role-based collaboration, and human checkpoints, with orchestration metadata captured in registries for traceability.

Cross-Cutting Requirements

Regardless of role or access model, Scaffolding ensures:

- ♦ Registries for all GPTs, agents, prompts, and tools, including ownership, purpose, and risk tier.
 - ♦ Central runtime enforcement and data protection through the Agent & MCP Gateway.
 - ♦ Fine-grained access control using RBAC/ABAC models.
 - ♦ Full observability across usage, risk, and cost dimensions.
- 

VII. Program Structure

Moving from principles to practice requires deliberate execution. A practical way to operationalize AI Scaffolding is through four complementary workstreams. These are not sequential phases but parallel, interdependent streams that reinforce each other and mature together:

- ◆ Assessment
- ◆ Policy Development
- ◆ Infrastructure Building
- ◆ Enforcement (via Guardian Agents)

Workstream A

- Assessment

(Guardrails & Environment Discovery)

Assessment establishes and maintains the baseline — surfacing current usage, risk exposure, and gaps — so Scaffolding can be prioritized effectively and adapted as new AI assets emerge.

Objective:

Establish a comprehensive view of the current AI environment, risk tiers, and existing guardrails.

Illustrative Deliverables:

- ◆ Guardrail taxonomy and severity thresholds.
- ◆ Baseline metrics and coverage map.
- ◆ Initial draft registries for GPTs, agents, prompts, and tools.
- ◆ Assessment report highlighting quick-win opportunities.

Workstream B

- Policy Development

(Governance, Metadata & Registries, Tool Standards)

Policy development defines the rules of the road. Enablement ensures the organization knows how to follow them and why they matter.

Objective:

Define governance frameworks, approval flows, and metadata standards that underpin safe and scalable AI adoption.



Illustrative Deliverables:

- ♦ AI Governance Charter & Policy Pack – including acceptable use, citizen developer enablement, prompt/data standards, and retention/logging.
- ♦ Registry Governance Policy & JSON Schemas – covering GPTs, agents, MCP tools, and prompts.
- ♦ Risk Register aligned to NIST AI RMF and ISO/IEC 42001.
- ♦ Interoperability Standards for Agent Cards and MCP connectors.

Workstream C

- Infrastructure Building

(Policy-as-Code, Registries, Gateway)

Infrastructure operationalizes Scaffolding - delivering the registries, gateways, and observability needed for safe scaling.

Objective:

Build the Scaffolding infrastructure that enforces policy and provides shared services once, centrally, across the enterprise.

Illustrative Deliverables:

- ♦ Running registries for GPTs, agents,

prompts, and MCP tools.

- ♦ Agent & MCP Gateway MVP with runtime policy enforcement and sensitive-data protection.
- ♦ Policy-as-Code Repository containing guardrails and evaluation pipelines.
- ♦ Observability Dashboards integrating OpenTelemetry, SIEM, and APM for centralized visibility.

Workstream D

- Enforcement

(via Guardian Agents)

Enforcement ensures that the rules defined in policy are actually followed in practice. While traditional enforcement relies on static gates and manual reviews, the modern approach uses Guardian Agents to embed enforcement directly into runtime workflows.

Guardian Agents also deliver continuous, non-human-in-the-loop (non-HITL) assurance – bringing automated structure to the lifecycle of AI assets - ensuring new agents, prompts, and tools are onboarded and evaluated before rollout, monitored continuously in production, and safely deprecated when retired.

Objective:

Operationalize policies through autonomous Guardian Agents that enforce controls before, during, and after execution.

Illustrative Deliverables:

- ♦ Guardian Agents as active enforcement components:
 - Gatekeeper – pre-invoke validation.
 - Guardian – inline content and risk filtering.
 - Auditor – post-hoc evaluation and drift detection.
- ♦ Registrar Agent enforcing registry compliance prior to go-live.
- ♦ Test Harness and Sampling Plan for quality assurance.
- ♦ Red-Team Playbook for proactive risk discovery.
- ♦ Lifecycle & Assurance Model – structured onboarding, rollout, and retirement processes, with continuous evaluation, versioning, red-teaming, and drift detection.



Strategic Outcomes

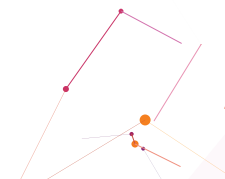
When executed through these workstreams, Scaffolding delivers outcomes aligned directly with enterprise priorities:

- ♦ **Visibility** – Continuous assessment and updated registries as new GPTs, agents, prompts, and tools emerge.
- ♦ **Consistency** – Governance frameworks mapped to standards such as **NIST AI RMF**, **ISO/IEC 42001**, and the **EU AI Act**.
- ♦ **Interoperability** – **Agent Cards** and **MCP connectors** that prevent fragmentation and vendor lock-in.

- ♦ **Adoption at Scale** – Defined personas, journeys, and workflows that make governance tangible for business users.
- ♦ **Shared Services** – Registries, gateways, policy-as-code, observability, and FinOps delivered once, centrally.
- ♦ **Autonomous Assurance** – Guardian Agents that enforce rules continuously across onboarding, rollout, monitoring, and retirement.

Clear Checkpoints — Structured 30/60/90-day decision gates with measurable KPIs, such as:

- ♦ % of assets registered and evaluated.
- ♦ SLA for approvals and exceptions.
- ♦ Registry completeness and runtime overhead.
- ♦ Incident detection rates, MTTR, and cost per guarded task.



VIII. Conclusion & Path Forward

The enterprises on the leading edge of AI are embracing full-speed adoption — a pace more typical of tech firms than of traditional industries. That boldness is already creating real value, putting them ahead of their peers not just in experimentation but in tangible outcomes.

Yet the wider business community is waiting for the first high-profile failure — the headline that proves boldness equals recklessness. Scaffolding is how enterprises can thread that needle: sustaining a revolutionary pace while making risk management tangible, automated, and proportionate.

With Scaffolding in place, companies don't have to choose between speed and safety. Governance shifts from a brake to an engine — the mechanism that allows adoption to keep accelerating securely and at scale. That is how true leaders will define themselves in enterprise AI.

Appendices

Governance & Risk Controls *(examples)*

Scaffolding consolidates governance and risk controls into shared services rather than leaving each team to solve them independently. Illustrative examples include:

Policy Gates

- ◆ Registration of models, GPTs, agents, prompts, and tools before production use.
- ◆ Risk tiering data sensitivity, use case criticality, and compliance exposure.
- ◆ Evaluation of thresholds and human-in-the-loop requirements for high-risk workflows.

Safety Controls

- ◆ Provider guardrails for toxicity prevention, jailbreak defense, and PII redaction.
- ◆ Retrieval grounding checks to prevent unsupported or hallucinated outputs.
- ◆ Red-teaming and adversarial testing processes.

Access & Permissions

- ◆ Least-privilege access for tools and agents, with auditable scopes and approvals.
- ◆ Consistent RBAC/ABAC enforcement across GPTs, agents, and MCP connectors.
- ◆ Secrets rotation and OAuth/OIDC authentication enforcement.

Change Management

- ◆ Versioning of prompts, agents, and tools.
- ◆ Rollback procedures and deprecation windows to ensure controlled evolution.

Monitoring & Observability

- ◆ Prompt and output logging with privacy redaction.
- ◆ Anomaly detection for cost, drift, and unusual usage patterns.
- ◆ OpenTelemetry and SIEM/APM integration for full traceability.

Incident Response:

- ◆ Runbooks for handling model misbehavior, hallucinations, data leakage, and access abuse.
- ◆ Escalation workflows and metrics (MTTR, severity levels).

These controls are well understood in other enterprise contexts (*identity, cloud, data*). Scaffolding applies to them consistently to the AI ecosystem, enabling speed with safety.

A. Experience & Adoption (examples)

Scaffolding is only effective if it translates into usable, trusted experiences for employees, builders, and governance teams. Illustrative adoption patterns include:

Personas

- ◆ Frontline Agent – Customer-facing; requires fast, accurate responses without exposing PII.
- ◆ Finance Analyst – Needs governed access to structured data and complete audit trails.
- ◆ Claims Adjuster – Benefits from workflow agents that integrate across multiple systems.
- ◆ Field Sales Representative – Seeks mobile AI assistance connected to approved data sources.

Request intake → policy validation → execution or agent action → human handoff if required → audit trail creation → continuous feedback for improvement.

Prototypes & Patterns

- ◆ Agent handoff between AI assistants and human operators.
- ◆ Multi-agent collaboration with clear role separation (e.g., triage vs. resolution).
- ◆ Human approval nodes for high-risk actions.

Citizen Developer Enablement

- ◆ Templates for prompt packs, tool stubs, and evaluation checks.
- ◆ Review lanes and publishing workflows that make it easy for employees to build safely.
- ◆ Guardrailed sandboxes that enable experimentation without risking sensitive data or systems.

Grounding Scaffolding in personas, journeys, and enablement flows ensures that governance is not just protective but also accelerates adoption by giving employees confidence to use AI tools in their day-to-day work.

B. Sample Agent Card (A2A) – Sample Agent Card (A2A)

An Agent Card provides a standardized way to describe an agent's identity, purpose, skills, and governance metadata. This enables discoverability in registries and interoperability across platforms.

Key Points Illustrated

- ◆ Clear ownership and purpose for every AI asset.
- ◆ Authentication and scope defined (least-privilege)
- ◆ Governance metadata captured including risk tier, evaluation status, and blast radius.
- ◆ Security groups and attributes explicitly defined for RBAC/ABAC enforcement.
- ◆ Interoperability ensured through MCP server references.
- ◆ Traceability achieved via metadata injection across workflows.

```
{
  "name": "Ops Helper",
  "description": "Resolves incidents and executes runbooks in ServiceNow",
  "version": "1.0.0",
  "owner": "IT Operations",
  "endpoint": "https://ops.example.com/a2a",
  "authentication": { "schemes": ["Bearer"] },
  "capabilities": { "streaming": true },
  "skills": [
    {
      "id": "create_incident",
      "name": "Create Incident",
      "description": "Open a Sev-2 incident in ServiceNow",
      "inputModes": ["application/json"],
      "outputModes": ["application/json"]
    }
  ],
  "mcp": {
    "defaultServer": "servicenow",
    "servers": [
      {
        "id": "servicenow",
        "url": "https://tools.example.com/mcp/servicenow",
        "auth": {
          "type": "oauth2",
          "client": "machine-to-machine",
          "scopes": ["incident:create", "incident:read"]
        }
      }
    ]
  },
  "governance": {
    "riskTier": "Moderate",
    "evalStatus": "Approved",
    "logging": "Full",
    "blastRadius": "Scoped to IT incidents only"
  },
  "securityGroups": [
    {
      "groupId": "ops_incident_handlers",
      "description": "Restricts access to authorized operations staff only",
      "attributes": ["RBAC:Incident:Write", "RBAC:Incident:Read"]
    }
  ],
  "metadataInjection": {
    "requiredClaims": ["user_id", "purpose", "context_id"],
    "propagationHeaders": ["x-correlation-id", "x-trace-id"]
  }
}
```

C. Appendix B – MCP Tool Registry

An MCP Tool Registry entry describes a tool's owner, purpose, data access, and governance attributes. By standardizing this metadata, enterprises ensure safe interoperability between agents and tools.

Key points illustrated here:

- ◆ Ownership & Purpose are explicit.
- ◆ Authentication, scopes, and RBAC/ABAC groups are defined.
- ◆ Risk-tiering signals sensitivity.
- ◆ Approval authority establishes accountability.
- ◆ Versioning and rollback policies support lifecycle management.
- ◆ Endpoints clarify whether the registry calls direct data sources or whitelisted APIs inheriting security attributes.
- ◆ Governance rules enforce change control and allow emergency disable.

Together, Agent Cards (*Appendix C*) and MCP Tool Registry entries (*Appendix D*) show how Scaffolding makes agents and tools discoverable, governable, and interoperable across the enterprise.

```
{
  "toolId": "servicenow_incident",
  "name": "ServiceNow Incident Tool",
  "description": "Create and retrieve IT incidents from ServiceNow",
  "owner": "IT Service Management",
  "purpose": "Incident creation and resolution support",
  "authModel": "OAuth2 Client Credentials",
  "scopes": ["incident:create", "incident:read"],
  "riskTier": "High",
  "dataClasses": ["Operational Data"],
  "approvalAuthority": "Enterprise IT Governance Committee",
  "loggingLevel": "Full",
  "retention": "90 days",
  "evaluationStatus": "Approved",
  "versioning": {
    "currentVersion": "1.2.0",
    "deprecationWindow": "6 months",
    "rollbackPolicy": "Previous stable version"
  },
  "securityGroups": [
    {
      "groupId": "ops_incident_handlers",
      "description": "Restricts tool usage to approved incident management teams",
      "attributes": ["RBAC:Incident:Read", "RBAC:Incident:Write"]
    }
  ],
  "governance": {
    "requiredMetadata": ["owner", "purpose", "scopes", "riskTier"],
    "changeRules": "New scopes require re-approval",
    "emergencyDisable": true
  },
  "endpoints": [
    {
      "type": "direct",
      "url": "https://api.servicenow.com/incidents",
      "inheritsSecurityGroups": false
    },
    {
      "type": "whitelistedAPI",
      "url": "https://api.mycompany.example.com/incident-proxy",
      "inheritsSecurityGroups": true
    }
  ]
}
```

D. Appendix E - Prompt Registry Entry

A Prompt Registry provides a structured way to manage prompt engineering assets - ensuring they are versioned, governed, and auditable just like agents and tools. Prompts can be treated as first-class artifacts with metadata around ownership, purpose, lineage, and evaluation status..

Key points illustrated here:

- ◆ Prompts are treated as governed, versioned assets.
- ◆ Ownership, purpose, and risk tier are explicit.
- ◆ Security groups limit who can deploy or invoke the prompt.
- ◆ Versioning, rollback, and deprecation policies ensure lifecycle control.
- ◆ Lineage metadata tracks training data sources and derived versions.
- ◆ Assurance metadata links to evaluation coverage, benchmarks, and red-team testing

```
{
  "promptId": "claims_intake_v3",
  "name": "Claims Intake Assistant Prompt",
  "description": "Collects and normalizes claim information from policyholders",
  "owner": "Claims Operations",
  "purpose": "Standardize intake data for downstream processing",
  "dataClasses": ["Customer PII", "Claims Data"],
  "riskTier": "High",
  "approvalAuthority": "AI Governance Committee",
  "evaluationStatus": "In Evaluation",
  "versioning": {
    "currentVersion": "3.0.1",
    "previousVersions": ["3.0.0", "2.5.2"],
    "rollbackPolicy": "Rollback permitted to last approved version",
    "deprecationWindow": "12 months"
  },
  "securityGroups": [
    {
      "groupId": "claims_intake_users",
      "description": "Restricts use to authenticated claims adjusters and intake agents",
      "attributes": ["RBAC:Claims:Write", "RBAC:Claims:Read"]
    }
  ],
  "governance": {
    "requiredMetadata": ["owner", "purpose", "riskTier", "evaluationStatus"],
    "changeRules": "Material prompt changes (structure, tone, data references) require re-approval",
    "emergencyDisable": true
  },
  "lineage": {
    "trainingDataSources": ["internal_claims_knowledgebase_v2"],
    "derivedFrom": ["claims_intake_v2"]
  },
  "assurance": {
    "evalCoverage": "80%",
    "safetyBenchmarks": ["BiasCheck", "PIIRedactionTest"],
    "redTeamStatus": "In Progress"
  }
}
```

About Altimetrik

Altimetrik is an AI-first data- and digital-engineering company that helps enterprises accelerate growth with an incremental, product-oriented approach. An official services partner of OpenAI, Altimetrik's ALTI AI Adoption Lab™ and DomainForge.ai help enterprises build and deploy enterprise-grade AI solutions. With more than 6,000 practitioners worldwide and deep engineering DNA, we enable organizations across BFSI, manufacturing, retail and CPG, automotive, health care, and life sciences to modernize technology, launch new business models, and scale AI adoption.

Recognized in the 2025 Constellation Research ShortList™ for Global AI Services and named a Major Contender in Everest Group's PEAK Matrix® for BFSI IT Services Specialists and Life Sciences Digital Engineering Services, Altimetrik ensures efficiency, visibility, and frictionless processes, empowering businesses to thrive in the AI era.



www.altimetrik.com