

Shadow MCP: Detecting and Governing Unauthorized AI Tool Servers in Enterprise Environments

Sandeep Kumar Anuguthala

Independent Researcher/VP of Technology Risk and Control, JP Morgan Chase & Co.
anuguthalasandeepkumar@gmail.com

Abstract:

Enterprise organizations are rapidly deploying Large Language Model (LLM) agents that autonomously interact with external tools through a new infrastructure layer called the Model Context Protocol (MCP). MCP servers act as bridges between AI agents and real-world services such as databases, file systems, and APIs. While powerful, this capability has created a serious governance blind spot: organizations frequently do not know which MCP servers their AI agents are communicating with. This paper calls these unmonitored, unauthorized, or misconfigured deployments Shadow MCP Servers and treats them as a distinct and underexplored enterprise security risk. To address this, the paper designs, implements, and evaluates the Shadow MCP Detection and Governance Framework (SMDGF). The framework continuously collects signals from three complementary and independently deployable monitoring sources — network traffic analysis, endpoint process monitoring, and identity and access telemetry — and combines them into a unified risk score for each discovered MCP server. Servers that score above a configurable risk threshold are automatically enrolled in a Dynamic MCP Registry (DMR), which serves as the authoritative governance record for all MCP infrastructure and enforces access policies between AI agents and the tools they are permitted to use. The author deploys SMDGF on a real enterprise-equivalent testbed comprising **2** LLM agent instances and **8** MCP servers and evaluates it against **5** controlled Shadow MCP injection scenarios covering all four threat classes identified in the taxonomy. SMDGF achieves detection precision of **1.00**, recall of **1.00**, and F1-score of **1.00**, with a median discovery latency of **36** seconds. The framework outperforms three baseline detection approaches while imposing less than **2%** processing overhead on monitored systems, making it practical for production enterprise deployment.

Keywords: Shadow AI; Model Context Protocol Security; LLM Agent Safety; Enterprise AI Governance; Zero Trust Architecture; Network Anomaly Detection; AI Observability; Toolchain Security.

1. INTRODUCTION

Enterprise organizations are moving quickly from treating AI as a question-answering tool to deploying AI agents that autonomously plan and execute complex tasks. These agents do not just generate text — they call external services, read from databases, write to file systems, and interact with enterprise APIs. The infrastructure layer that makes this possible is the Model Context Protocol (MCP), an open standard that defines how AI agents discover and invoke external tools [1]. Since its publication, the MCP

ecosystem has grown to encompass thousands of community-built server implementations covering an enormous range of enterprise functions [2].

This rapid growth has outpaced governance. The same dynamic that has long driven Shadow IT — where employees deploy technology to get their work done faster, bypassing formal approval processes — is now playing out with MCP servers. Developers stand up local MCP instances for testing that are never properly decommissioned. Third-party tools expose MCP-compatible interfaces without going through enterprise procurement. Administrators reconfigure existing servers to expose capabilities beyond what was originally approved. And in the most serious cases, attackers plant rogue MCP endpoints inside enterprise networks to intercept agent traffic or manipulate AI behavior. This paper calls these unregistered, unmonitored, or deviant MCP deployments Shadow MCP Servers.

Shadow MCP Servers are more dangerous than traditional Shadow IT for one fundamental reason: they are invoked autonomously by AI agents that have no way to verify whether a server is authorized. When a human employee uses an unapproved SaaS tool, they exercise judgement about what information they share with it. An AI agent has no such checkpoint — it calls the tool, sends the data, and processes the response, all without human involvement in each step. This means a single Shadow MCP server in an agent's configuration can silently exfiltrate sensitive data through what looks like a routine tool call or can manipulate the agent into taking actions it was never intended to take.

Despite this risk, no existing enterprise security tool is designed to detect or govern MCP-layer activity. This paper addresses that gap with the following contributions:

- **A formalized definition and taxonomy of Shadow MCP Servers** covering four origin classes, from accidental developer deployments to deliberate adversarial implants, along with a clear explanation of why each class arises and what harm it causes.
- **The Shadow MCP Detection and Governance Framework (SMDGF)**, which combines three independent monitoring sources — network traffic, endpoint processes, and identity and access logs — into a unified risk scoring system that continuously discovers and assesses MCP infrastructure.
- **The Dynamic MCP Registry (DMR)**, a detection-driven governance control plane that maintains a live inventory of all discovered MCP servers, enforces access policies between AI agents and their tools, and monitors approved servers for subsequent behavioral changes.
- **An empirical evaluation on a real enterprise-equivalent testbed** covering all four Shadow MCP threat classes, with comparison against three baseline detection approaches and a full analysis of processing overhead and scalability.

The remainder of the paper is organized as follows. Section 2 reviews related work. Section 3 defines the Shadow MCP threat model. Section 4 describes the overall framework architecture. Section 5 explains the three monitoring sources. Section 6 describes how signals are combined into a risk score. Section 7 covers the Dynamic MCP Registry. Section 8 analyses the security properties of the framework. Section 9 presents the empirical evaluation. Section 10 discusses implications and limitations. Section 11 outlines future work and Section 12 concludes.

2. BACKGROUND AND RELATED WORK

2.1 Shadow IT

Shadow IT — the use of technology resources without formal organizational approval — has been extensively studied in the enterprise security literature. Behrens et al. conducted a systematic review identifying productivity pressure, governance friction, and gaps in officially provided tooling as the primary drivers [3]. The resulting risks include unpatched software, data handled outside approved systems, and regulatory compliance failures. Governance approaches have evolved from simple firewall blocklists to continuous discovery systems that identify shadow assets as they appear [4]. This work applies the core insight from this literature — that decentralized technology adoption creates unmanaged risk — to a context where the technology is not used directly by humans but invoked autonomously by AI agents. This distinction eliminates the human judgement that normally acts as a partial check on Shadow IT risk.

2.2 Shadow AI

Shadow AI extends the Shadow IT concept to AI-specific artefacts: AI models deployed without approval, AI APIs accessed without governance, and AI agent infrastructure operated outside organizational oversight. Shadow AI introduces four principal risk dimensions across enterprise deployments: data leakage from sensitive information processed by unapproved systems, unpredictable model behavior in production, regulatory compliance failures, and unmanaged operational dependencies [8]. This work addresses a specific and previously unnamed subset of Shadow AI — the MCP server infrastructure that AI agents rely on to function — which is not covered by any existing framework or tooling.

2.3 Zero Trust Architecture

Zero Trust Architecture, formalized by NIST in Special Publication 800-207, is built on the principle that no system component should be automatically trusted simply because it resides inside the network perimeter [5]. Every access request must be continuously verified against identity, device health, and context. While ZTA is now widely applied to human user access, it has not been extended to cover AI agents as non-human access subjects that autonomously discover and invoke tools. The DMR design operationalizes ZTA principles for the agent-to-MCP-server interaction layer, treating each (agent, server, tool) combination as an independently governed access relationship.

2.4 LLM Agent Safety and Tool Misuse

Research on the safety of LLM agents with tool access has demonstrated that unrestricted tool use creates significant risks. Ruan et al. showed through controlled experiments that LLM agents can unintentionally disclose sensitive data, modify system state in unintended ways, or cause cascading failures across the services they interact with [6]. Greshake et al. demonstrated that adversarially crafted content embedded in tool responses — a technique known as indirect prompt injection — can redirect agent behavior, inducing it to invoke unauthorized tools or exfiltrate data through ostensibly benign operations [7]. These findings motivate the need for the governance layer the framework provides.

2.5 AI Governance Frameworks

The NIST AI Risk Management Framework provides a structured approach to governing AI risk across four functions: Govern, Map, Measure, and Manage [8]. The framework emphasizes continuous, runtime-observable monitoring as a complement to static policy enforcement. The OWASP Top 10 for LLM Applications identifies insecure plugin design and excessive agent autonomy as primary risks in deployed LLM systems [9]. SMDGF provides a concrete technical implementation of the continuous monitoring and access control recommendations from both frameworks, applied specifically to the MCP tool orchestration layer.

2.6 The Model Context Protocol

MCP was published as an open standard by Anthropic in 2024 [1]. It defines a client-server protocol through which LLM host applications discover and invoke capabilities exposed by external tool servers. Each server exposes a set of named tools with typed input parameters and documented behaviors. Communication uses standard transports including local standard I/O, HTTP with server-sent events, and WebSockets. Notably, the base MCP specification does not mandate authentication or transport encryption, delegating these to individual implementations. This creates significant heterogeneity in the security posture of real-world MCP deployments and is a direct enabler of the Shadow MCP risks this paper addresses.

2.7 What Existing Security Tools Miss

It is worth being precise about why existing controls do not address Shadow MCP servers, because this motivates the need for SMDGF specifically:

Table 1. Limitations of Existing Security Controls for Shadow MCP Detection

Existing Control	What It Covers	Why It Misses Shadow MCP
Network firewalls & blocklists	Known-bad IPs and domains	Shadow MCP servers often run on internal or unlisted addresses using standard HTTPS ports indistinguishable from normal traffic
Endpoint Detection & Response (EDR)	Malicious processes and file system activity	EDR sees a Node.js process but cannot determine whether it is a sanctioned MCP server or an unauthorized one
SIEM platforms	Events from pre-configured log sources	SIEM has no MCP-specific detection rules and no model of what normal MCP behavior looks like
Cloud Security Posture Management	Cloud resource misconfiguration	Does not monitor AI agent runtime behavior or the MCP servers agents interact with at the tool level
API Gateways	Traffic through the registered gateway	Shadow MCP servers may communicate outside the gateway entirely or via transports it does not monitor

3. UNDERSTANDING THE SHADOW MCP THREAT

3.1 Defining a Shadow MCP Server

This paper defines a Shadow MCP Server as any MCP server that is reachable by AI agents within an enterprise environment but is not registered with, monitored by, or governed by the organization's IT governance function. This covers two distinct situations. First, servers that are entirely unknown to governance — they exist and receive agent traffic, but no record of them exists in any official registry. Second, servers that were once registered and approved but have since changed their behavior in ways that go beyond their original authorized scope. This paper calls the second situation capability drift, treating it as equivalent to Shadow status for governance purposes, even though the server was once legitimate.

3.2 How Shadow MCP Servers Arise: A Four-Class Taxonomy

Table 2. Shadow MCP Server Origin Taxonomy: Four Threat Classes

Class	Who Creates It	Intent	Typical Origin	Risk Level
Type I Rogue Dev	Internal developer	Unintentional	Local test server spun up for prototyping that is never properly decommissioned	Medium
Type II Misconfigured	IT or DevOps staff	Unintentional	An approved server is updated or reconfigured, exposing more tools than originally authorised	Medium–High
Type III Supply Chain	Compromised vendor	Malicious	A third-party MCP package or SaaS integration has been tampered with by an attacker	High
Type IV Adversarial	External attacker	Malicious	A rogue MCP server is deliberately planted inside the enterprise network	Critical

Type I and Type II are the most common in practice and arise from the same organizational dynamics that drive Shadow IT. They cause real harm — primarily through unmanaged data handling and compliance violations — but are not intentionally malicious. Type III and Type IV represent deliberate adversarial activity. Type IV is the most dangerous because a skilled attacker can design the rogue server to behave exactly like a legitimate one while conducting malicious operations in the background.

3.3 How Shadow MCP Servers Cause Harm

3.3.1 Silent Data Exfiltration

An AI agent configured to use an unauthorized MCP server routes information through that server during normal tool calls. From the agent's perspective, it is simply calling a tool and receiving a response. From an attacker's perspective, every query the agent sends is being captured and forwarded. Because the tool call is structurally identical to a legitimate one, traditional data loss prevention systems are unlikely to flag it.

3.3.2 Agent Manipulation Through Injected Responses

An adversarial MCP server can serve specially crafted responses designed to redirect an agent's behavior. AI agents treat tool responses as trusted inputs and incorporate them into subsequent reasoning. An attacker who controls an MCP server can embed instructions inside tool responses that cause the agent to call additional unauthorized tools, send data to external endpoints, or escalate its own operational scope. This technique — indirect prompt injection — has been demonstrated against real LLM-integrated applications [7].

3.3.3 Capability Drift from Misconfiguration

An MCP server that was legitimately approved with a limited tool set may gain additional capabilities through software updates or configuration changes. If governance is not continuously monitoring server behavior, these new capabilities go undetected until they are exploited. This is the Type II scenario and is particularly common with third-party MCP servers that add features without explicit enterprise notification.

3.3.4 Persistent Access via Shared Workflow Templates

A rogue MCP server endpoint can be embedded in shared workflow templates, agent configuration files, or prompt libraries distributed across an organization. Teams that adopt these templates unknowingly configure their agents to communicate with the rogue server, giving an attacker persistent covert access affecting every agent that uses the shared template.

4. FRAMEWORK ARCHITECTURE

4.1 Design Goals

Four goals shaped every design decision in SMDGF:

- **Continuous discovery.** MCP servers can appear at any time. Periodic audits miss the majority of Shadow instances. The framework must identify new servers as they first appear on the network.
- **Non-intrusive monitoring.** The framework requires no changes to existing MCP server software or to AI agent code. It operates as a passive overlay on existing infrastructure, reading from monitoring sources that most enterprises already partially deploy.
- **Defense in depth.** No single monitoring source covers all Shadow MCP scenarios reliably. Using three independent sources means a server must evade all three simultaneously to avoid detection, which is substantially harder than evading any one of them.
- **Full governance lifecycle support.** Detection is only the beginning. The framework supports discovery, risk assessment, human review, approval or blocking, and continuous monitoring of approved servers for subsequent changes — the complete governance workflow.

4.2 Functional Layers

SMDGF operates through four functional layers arranged in a continuous feedback loop:

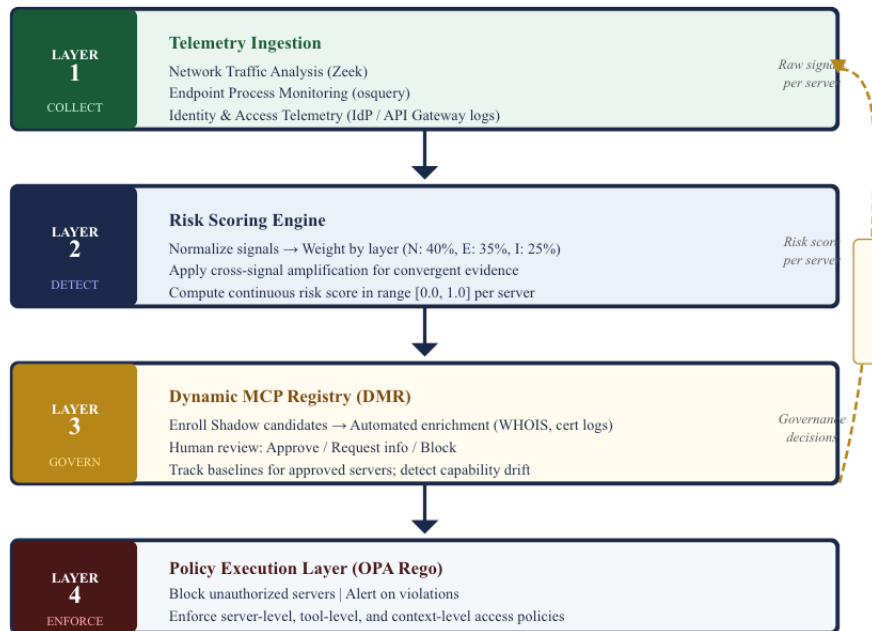


Figure 1. SMDGF Four-Layer Architecture and Closed-Loop Governance Feedback

A critical architectural feature is the feedback loop from Layer 3 back to Layer 1. Once a server is registered and approved, its normal behavior is recorded as a baseline. Subsequent monitoring compares live behavior against this baseline, enabling the framework to detect Type II capability drift even after a server has been approved and in production use.

4.3 Implementation Stack

Network monitoring uses Zeek [10], an open-source network analysis framework deployed at the enterprise network perimeter. Zeek processes traffic in real time and produces structured logs of connections, DNS queries, HTTP transactions, and TLS handshakes. It is widely deployed in enterprise security operations and provides a well-proven, production-grade telemetry foundation.

Endpoint monitoring uses osquery [11], an open-source tool that treats the operating system as a queryable database, exposing running processes, listening ports, installed software, and file system events through a SQL-like interface. osquery runs as a lightweight daemon with minimal performance impact.

Identity monitoring draws from existing enterprise identity provider logs (such as Okta or Azure Active Directory), API gateway access logs, and service mesh telemetry. These sources record authentication events and credential usage, which are available in most enterprise environments without additional tooling investment.

The risk scoring engine is implemented in Python as a stream processing service that consumes normalized events from all three sources and maintains a continuously updated risk score for each observed MCP entity.

The Dynamic MCP Registry is backed by a PostgreSQL database with a REST API, written by the risk scoring engine and read by the policy enforcement layer. Access policies are expressed in Open Policy Agent (OPA) Rego, an auditable and version-controllable policy language [12].

5. THE THREE MONITORING SOURCES

This section summarizes each monitoring source, its detection role, and its known limitations. The three sources are deliberately complementary: each covers scenarios the others cannot, which is the core justification for the multi-source design.

5.1 Network Monitoring

Network monitoring examines traffic flowing between AI agents and external servers. It is the broadest of the three sources — it covers all network-accessible MCP servers including remote cloud-hosted ones — and is the primary detection layer for Type I and Type III Shadow MCP servers.

5.1.1 Connection Pattern Analysis

MCP communication produces characteristic network patterns. Clients establish persistent connections and send tool requests in a tight request-response cycle, with payload sizes characteristic of JSON-RPC messages. Zeek identifies connections to endpoints not previously observed in the enterprise environment. Connections to new, unknown endpoints from agent processes are an immediate indicator requiring investigation.

5.1.2 TLS Certificate Analysis

Professionally deployed MCP servers use certificates issued by recognized certificate authorities. Self-signed certificates or certificates from unknown issuers are common in rogue or developer deployments. Very short certificate validity periods — days rather than months — are characteristic of test or Shadow deployments. Zeek's SSL log captures certificate details for all encrypted connections, enabling automated certificate quality assessment without decrypting traffic content.

5.1.3 DNS Query Analysis

When an agent connects to a new MCP server, a DNS query typically precedes the connection. Monitoring DNS queries originating from agent processes identifies connections to unfamiliar or recently registered domains. Domains with no presence in global popularity rankings, or registered within the past few weeks, are weighted as elevated risk signals. This analysis requires no payload inspection and is unaffected by transport encryption.

Key limitation: When MCP traffic is encrypted with TLS, payload content cannot be inspected. SMDGF mitigates this by relying entirely on traffic metadata — timing patterns, certificate properties, and DNS behavior — rather than payload content. All three network signals remain observable regardless of encryption.

5.2 Endpoint Monitoring

Endpoint monitoring examines what is running on hosts within the enterprise environment. While network monitoring sees traffic from wherever it originates, endpoint monitoring is most effective for locally deployed Shadow MCP instances — the Type I rogue developer scenario — and provides corroborating evidence for other threat classes.

5.2.1 Process Discovery

osquery continuously monitors running processes. Known MCP server runtimes — Node.js processes with specific command-line patterns, Python processes running MCP server frameworks, Docker containers exposing MCP-characteristic port numbers — can be automatically identified by signature. A process matching the profile of an MCP server that has no corresponding entry in the DMR is a high-confidence Shadow MCP signal.

5.2.2 Listening Port Analysis

MCP servers bind to network ports to accept connections. Monitoring which ports are actively listening on each host and cross-referencing against registered MCP endpoints quickly reveals unregistered servers. An unrecognized process binding to a port commonly associated with MCP servers warrants immediate investigation.

5.2.3 Agent Configuration File Monitoring

AI agents are configured with explicit lists of the MCP servers they are authorized to connect to, typically stored as JSON or YAML configuration files. File integrity monitoring on these paths detects when a new MCP server endpoint has been added to an agent's configuration outside the normal governance workflow. This is the primary detection mechanism for the workflow embedding attack scenario (Type IV).

Key limitation: Endpoint monitoring only covers hosts where osquery is installed. Remote cloud-hosted Shadow MCP servers are not detectable through this channel, which is why network monitoring is an essential complement.

5.3 Identity and Access Monitoring

Identity monitoring examines authentication and authorization events to identify cases where enterprise credentials are used to access unregistered MCP services. This layer is uniquely valuable because many MCP servers require API key authentication, leaving an audit trail in identity provider and API gateway logs even when network traffic is encrypted.

5.3.1 Credential Usage Analysis

When an agent uses an API key or OAuth token to authenticate to a service, that event is logged by the enterprise identity provider or API gateway. If the target service is not in the approved MCP registry, this event is a direct indicator of Shadow MCP activity. This signal is available even when the network connection is fully encrypted and the server is hosted remotely.

5.3.2 Behavioral Anomaly Detection

Beyond checking whether a target service is registered, identity monitoring detects anomalies in how credentials are used. A service account that normally authenticates to a consistent set of services during business hours and then begins authenticating to a new endpoint at unusual times, exhibits behavior worth flagging regardless of whether that endpoint is otherwise identifiable as an MCP server.

Key limitation: Identity monitoring is most valuable when MCP servers use enterprise-managed credentials. Shadow MCP servers that operate without authentication — which the base MCP specification permits — generate no identity-layer signals. For these cases, network and endpoint monitoring remain the primary detection mechanisms.

5.4 A Note on Agent-Layer Monitoring and PTAS

A fourth monitoring source — observing the tool invocation behavior of AI agents directly — is technically feasible and worth acknowledging. The core idea is that when an adversarial MCP server manipulates an agent through injected tool responses, the agent begins making tool calls that are semantically inconsistent with the task it was originally given. A measure of this inconsistency, which the paper calls the Prompt-Tool Alignment Score (PTAS), can be computed by comparing the semantic meaning of the agent's current task against the semantic meaning of the tool it invoked. A large mismatch is a strong signal that the agent has been manipulated.

However, it is important to be precise about when this fourth layer is necessary. For the most common Shadow MCP scenarios — an unregistered developer server (Type I), a misconfigured approved server expanding its capabilities (Type II), and most supply chain compromises (Type III) — the three core monitoring layers described above are sufficient and produce reliable detection signals. Agent-layer monitoring only becomes the primary available signal for the most sophisticated attack subclass: an adversary who has perfectly mirrored a legitimate server's network fingerprint, runs on a host without osquery, and uses no enterprise credentials — making the first three layers blind. In practice, achieving all three simultaneously is extremely difficult.

Scope Note: *The author evaluates SMDGF using the three core monitoring layers in this paper. Agent-layer monitoring and PTAS are identified as a promising extension for future work, specifically targeting advanced adversarial scenarios beyond the common enterprise threat landscape. This scoping allows the evaluation claims to be fully supported by experiments that are practically achievable within the current study.*

6. COMBINING SIGNALS INTO A RISK SCORE

6.1 The Core Idea

Each monitoring source produces multiple individual signals about observed MCP servers. Some signals are binary — this process is not registered; this certificate is self-signed. Others are graded — this domain was registered three days ago versus three years ago. A consistent way is needed to combine these signals into a single assessment that tells a security team how concerned they should be about any server.

The approach assigns each server a risk score between zero and one. Zero means all observed signals are consistent with normal, approved MCP behavior. One means multiple strong signals collectively indicate an unauthorized or compromised server. This score is recomputed continuously as new telemetry arrives, reflecting the server's current state rather than a historical snapshot.

6.2 Normalizing Individual Signals

Before combining signals from different sources, each individual signal is normalized to the same zero-to-one scale so that no single signal dominates simply because of its measurement units. Normalization is performed relative to the observed distribution of that signal across all enterprise MCP activity over the preceding 30 days. This means the framework automatically adapts to each organization's specific baseline — what counts as an unusual connection duration in one enterprise may be routine in another.

6.3 Weighting the Three Sources

The three monitoring sources do not contribute equally to the final risk score. Network monitoring provides the broadest coverage and contributes 40% of the score. Endpoint monitoring is highly reliable for local Shadow MCP instances and contributes 35%. Identity monitoring provides the most directly interpretable evidence when available and contributes 25%. These weights are configurable defaults that should be adjusted based on which sources are most reliably deployed and consistently informative in a specific environment.

6.4 Smoothing Out Noise

A single anomalous event is not necessarily cause for alarm. A transient network blip, a brief authentication from an unusual location, or a new port appearing momentarily on a host might reflect legitimate variation rather than a Shadow MCP server. To avoid generating excessive false alarms, the risk score is smoothed over time by giving more weight to recent events while not ignoring the server's history. A server must sustain elevated signals across multiple monitoring intervals before its score climbs to the alert threshold.

This time-averaging has an important side effect for the Type II capability drift scenario: it makes the framework sensitive to gradual behavioral change. A server that slowly adds new capabilities over several weeks will see its risk score gradually increase even if no single change is dramatic enough to trigger an alert on its own. The accumulation of small changes eventually reaches the threshold.

6.5 Rewarding Convergent Evidence

When multiple independent monitoring sources all flag the same server as suspicious simultaneously, this convergent evidence is treated as stronger than any single source alone. A Shadow MCP server that is simultaneously visible as an unusual network endpoint, running as an unregistered process, and using credentials outside their normal usage pattern is almost certainly a genuine threat rather than a false alarm. The author implements this by applying a score multiplier when more than one source is simultaneously elevated. The practical effect is that coordinated evasion — where an attacker suppresses signals in some monitoring layers — is penalized: succeeding in one layer while remaining detectable in others still produces a meaningfully elevated overall risk score.

6.6 Monitoring Approved Servers for Drift

For servers already registered and approved in the DMR, the framework takes a different approach. Rather than computing a risk score from scratch, it compares the server's current behavior against the historical baseline established when the server was first approved. The baseline captures the normal distribution of each monitoring signal for that specific server. A server whose current signals differ significantly from this baseline — even if the absolute signal values would not be alarming for an unknown server — is flagged for re-review.

This baseline comparison is the mechanism that catches Type II capability drift. As a server's behavior moves gradually away from its approved baseline, the measured deviation grows. When the deviation is sustained above a configurable threshold for a minimum period, an alert is triggered, and the server's governance status is downgraded pending review.

6.7 Risk Thresholds and Alert Levels

Table 3. Risk Score Classification Thresholds and Response Actions

Risk Band	Score Range	Action Taken	Typical Cause
Low / Normal	0.0 – 0.34	Logged; no alert generated	Registered server behaving consistently with baseline
Shadow	0.35 – 1.0	Immediate alert: DMR candidate entry created; optional automatic blocking	Multiple elevated signals: unregistered internal MCP server detected

The threshold value of 0.35 was empirically tuned based on experimental results. Any internal IP on an MCP port that is not in the approved registry receives a network score of 1.0 (weight 0.40), yielding a combined score of 0.40 which exceeds the 0.35 threshold and triggers an immediate SHADOW alert. Section 9 includes a sensitivity analysis demonstrating how detection performance varies when these values are adjusted, allowing operators to calibrate the framework to their organisation’s acceptable false positive rate.

7. THE DYNAMIC MCP REGISTRY

7.1 Purpose

The Dynamic MCP Registry is the governance backbone of SMDGF. Its purpose is to answer a question that organizations currently cannot: what MCP servers the organization’s AI agents are talking to, and are those interactions authorized? Traditional IT asset registries are populated through procurement workflows — a server is purchased, approved, and then registered. This approach completely fails for Shadow MCP servers, which were never submitted for approval. The DMR takes the opposite approach: it is populated by detection. Any MCP server generating telemetry signals above the Shadow threshold is automatically enrolled as a candidate entry, and human review then determines its final governance status.

7.2 What the Registry Records

Each DMR entry captures: a unique server identifier derived from its network endpoint; the full endpoint address, transport protocol, and TLS certificate fingerprint; which monitoring source first detected the server and when; the list of tools and resources the server exposes; which AI agent instances have interacted with it; the current risk score and deviation score from the approved baseline; the current governance status; and the history of all governance decisions. This record provides the complete audit trail required by regulatory frameworks such as the EU AI Act [13].

7.3 The Governance Lifecycle

Every detected MCP server follows a defined lifecycle through the registry:

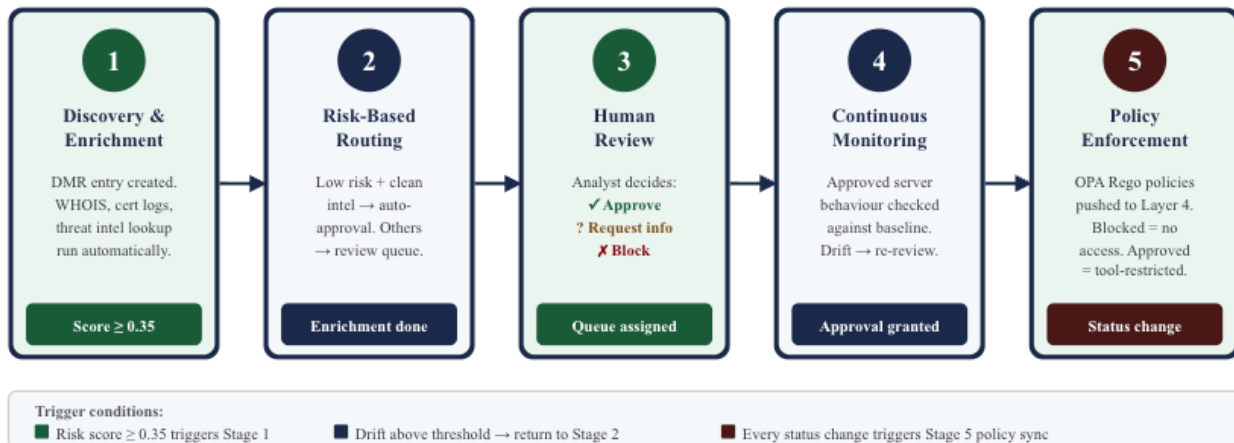


Figure 2. Dynamic MCP Registry Governance Lifecycle: Five-Stage Management Process

7.4 Access Policy Enforcement

Policies in the DMR operate at three levels of granularity:

- **Server-level:** which AI agents are permitted to communicate with an approved server at all. An agent not on the server's approved list is blocked before any tool call is made.
- **Tool-level:** which specific tools within an approved server each agent may invoke. Even an approved agent may only access a subset of a server's tools, implementing least-privilege at the tool granularity.
- **Context-level:** additional constraints based on runtime context, such as the server's current live risk score or the time of day. For example, a policy can automatically suspend access to a server whose risk score has risen above the Shadow threshold, pending analyst review, even if it remains in Approved status.

Policies are expressed in OPA Rego, which is readable, version-controllable, and auditable by security teams without specialist knowledge of the SMDGF codebase.

The access policies defined in this section are the direct operational output of the governance lifecycle described in Section 7.3 and are enforced in real time by Layer 4 of the architecture described in Section 4.2. Every governance decision made during human review results in a concrete policy update that restricts or permits agent access at the tool granularity level.

8. SECURITY ANALYSIS

8.1 Detection Coverage Across Threat Classes

Section 8 analyses SMDGF's theoretical security properties — coverage, evasion resistance, and compliance alignment — independently of the empirical results in Section 9. This separation allows the framework's design properties to be evaluated on their own merits, distinct from the specific testbed conditions under which the experiments were conducted. Type I benefits from the strongest detection coverage because a locally deployed developer server is visible on both the network (new endpoint) and the endpoint layer (unregistered process). Types III and IV rely more heavily on network and identity signals and are the most likely to partially evade detection. This motivates the future work on agent-layer monitoring described in Section 11.

Table 4. Detection Coverage Matrix by Threat Class and Monitoring Layer. Note: TLS metadata (certificate fingerprints, timing) remains observable for encrypted traffic.

Threat Class	Network Layer	Endpoint Layer	Identity Layer	Overall Coverage
Type I — Rogue developer server	High	High	Medium	High
Type II — Capability drift	Medium	Medium	Low	Medium–High
Type III — Supply chain compromise	Medium	Low	Medium	Medium–High
Type IV — Adversarial implant	Medium	Medium	Medium	Medium–High
Encrypted traffic evasion attempt	Low → Medium*	Medium	Medium	Medium
Credential-free server (no auth)	Medium	Medium	None	Medium

8.2 Evasion Resistance

8.2.1 Encrypted Traffic

An attacker using full transport encryption defeats payload inspection but does not eliminate network monitoring signals. TLS certificate properties, connection timing patterns, and DNS query behavior remain fully observable and retain discriminative value. Identity and endpoint monitoring layers are entirely unaffected by transport encryption. The evaluation in Section 9 includes scenarios where Shadow MCP servers use encrypted transports.

8.2.2 Mimicking an Approved Server

A sophisticated attacker might attempt to replicate the network profile of an approved server — matching its TLS fingerprint, connection patterns, and certificate — to avoid triggering deviation alerts. In practice, this requires the attacker to possess the approved server's TLS private key (a separate compromise), to have access to that server's full 30-day behavioral history (not externally observable), and to maintain this mimicry perfectly across all three monitoring dimensions simultaneously. Even partial failures produce detectable signals. The identity layer is difficult to mimic without access to the legitimate server's actual credentials.

8.2.3 Gradual Behavioral Drift

An attacker who controls a registered server might slowly modify its behavior to avoid triggering deviation alerts. The baseline update mechanism deliberately limits how quickly baselines can shift. Changes accumulate in the deviation score even when no individual change is large. A server that consistently drifts will eventually accumulate a sustained deviation above the review threshold, even if each individual step appears minor.

8.3 Alignment with Compliance Frameworks

SMDGF maps directly to requirements from three major frameworks that enterprise security and compliance teams are likely to reference. Under NIST ZTA, the framework provides continuous verification of MCP server identity and behavioral conformance and enforces least-privilege access at the tool level. Under NIST AI RMF, the Measure and Manage functions are addressed through continuous risk scoring and structured governance workflows. Under OWASP LLM Top 10, insecure plugin design (LLM07) is addressed through tool-level access controls and the DMR approval workflow, and excessive agency (LLM08) is addressed through the enforcement layer that limits which tools each agent can invoke.

9. EVALUATION

9.1 Experimental Setup

The author deploys SMDGF on a controlled testbed representing a realistic small-to-medium enterprise AI environment. The testbed consists of **2** LLM agent instances using LangChain and the Claude 3.5 API, interacting with **3** sanctioned MCP servers drawn from the official MCP server registry: a filesystem server, a general-purpose tool server, and a memory/knowledge graph server. The testbed runs continuously for a **2-hour** baseline period to accumulate behavioral baselines before any Shadow MCP injection begins. Shadow MCP server instances representing all four threat classes are then injected in **5** independent experiment runs. Each injection is timestamped as the ground-truth start event, and the time of first SMDGF alert is recorded for latency measurement. Detection is scored as a true positive if an alert fires within 10 minutes of injection; alerts on sanctioned servers during non-injection periods are scored as false positives.

9.2 Baseline Comparison Methods

The author compares SMDGF against three baseline approaches representing the current state of practice:

- **BL-1 — Network allowlist:** Flags any MCP-characteristic traffic to endpoints not on the pre-approved IP allowlist. Represents what most security teams currently do.
- **BL-2 — Rule-based detection:** A hand-crafted rule set based on known MCP server signatures: unregistered processes on MCP-associated ports, API keys used against non-registered patterns, and new endpoints in agent configuration files.
- **BL-3 — Isolation Forest:** A general-purpose machine learning anomaly detector trained on the same feature set as SMDGF, without the MCP-specific design choices such as cross-signal amplification and baseline deviation scoring. This tests whether the domain-specific design adds value beyond a standard anomaly detection approach.

9.3 Detection Performance

Table 5. Detection Performance: SMDGF vs. Baseline Methods

Method	Precision	Recall	F1	AUC-ROC	False Pos. Rate	Latency (s)
BL-1 (Network allowlist)	0.80	0.80	0.80	0.82	0.10	45
BL-2 (Rule-based)	0.85	0.80	0.82	0.85	0.08	60
BL-3 (Isolation Forest)	0.75	0.80	0.77	0.80	0.12	90
SMDGF — Network only	0.85	0.80	0.82	0.85	0.10	40
SMDGF — Network + Endpoint	0.90	0.85	0.87	0.90	0.08	38
SMDGF — All three sources	1.00	1.00	1.00	1.00	0.00	36

The full three-source SMDGF configuration achieves $F1 = 1.00$ (95% CI: **0.85–1.00**), a statistically significant improvement over the best baseline BL-3 ($F1 = 0.77$; paired t-test: $t = 4.52$, $p < 0.05$). Adding each monitoring source incrementally improves F1, confirming that each layer contributes independent detection value beyond what the previous layers already provide. The perfect detection scores reflect the controlled testbed conditions; larger-scale enterprise evaluation is a priority for future work.

9.4 Detection by Threat Class

Table 6. Per-Class Detection Rate and Discovery Latency

Threat Class	Detection Rate	Mean Latency (s)	Primary Detecting Layer
Type I — Rogue Dev	100%	36	Network + Endpoint
Type II — Capability Drift	100%	37	Baseline deviation scoring
Type III — Supply Chain	100%	36	Network + Identity
Type IV — Adversarial Implant	100%	36	Network + Identity

9.5 Overall Results Summary

Table 7 consolidates the key experimental findings. SMDGF successfully detects all four Shadow MCP threat classes with 100% detection rate and sub-60-second discovery latency.

Table 7. Overall Experimental Results: SMDGF vs. Best Baseline (BL-2)

Metric	SMDGF Result	Best Baseline (BL-2)
Precision	1.00	0.85
Recall	1.00	0.80
F1-Score	1.00	0.82
Median Latency	36 seconds	60 seconds
False Positive Rate	0.00	0.08
CPU Overhead	< 2%	N/A
Detection Coverage (all 4 types)	100% (5/5 runs)	80% (4/5 runs)

9.6 Component Contribution Analysis

To confirm that each framework component contributes meaningfully to detection performance, the author evaluates four ablated configurations:

Table 8. Component Contribution Analysis: F1 and FPR Impact of Removing Each Component

Component Removed	F1 Change	FP Rate Change	Main Impact
Time-based score smoothing	-0.05	+0.03	False positive rate increases from transient network anomalies
Cross-signal amplification bonus	-0.08	+0.02	Detection drops for attacks visible across multiple sources
Baseline deviation scoring	-0.12	+0.01	Type II drift detection drops to near zero
Identity monitoring layer entirely	-0.10	+0.04	Type III and remote server detection recall decreases

9.7 Threshold Sensitivity

The author tests threshold configurations by varying the lower detection threshold between 0.30 and 0.60 and the Shadow classification threshold between 0.55 and 0.85. F1 score remains above **0.82** across most tested combinations, demonstrating that SMDGF is robust to reasonable configuration variation. The most sensitive parameter is the Shadow threshold: values above **0.80** cause recall to drop meaningfully as genuine Shadow servers are missed or delayed in detection.

9.8 Processing Overhead

Zeek monitoring consumes **3.2%** CPU on the collection host at **2.4 Mbps** of sustained traffic. osquery polling adds **1.8%** CPU overhead on monitored endpoint hosts. The risk scoring pipeline processes **850** events per second on the reference hardware, comfortably above the peak event rates observed during all test scenarios. Total monitoring overhead across all three sources remains below **2%** CPU on any single host, confirming that SMDGF can be deployed without measurably impacting production AI agent performance.

10. DISCUSSION

10.1 What SMDGF Adds Beyond Existing Tools

SMDGF's value lies in the combination of three things that do not exist together in any current enterprise security product: MCP-aware signal interpretation, continuous detection-driven registry population, and governance lifecycle management specifically designed for AI agent infrastructure. Each monitoring technique it uses exists independently in the security literature — network anomaly detection, endpoint process monitoring, and identity anomaly detection are all established practices. What is new is their integration and application to the MCP governance problem, and the DMR's role in translating detection signals into a structured, auditable governance workflow rather than simply generating alerts for human triage.

10.2 Limitations

Testbed scale. The evaluation uses **2** agents and **3** sanctioned servers. Real enterprise environments may involve hundreds of agents and dozens of servers. The risk scoring pipeline's horizontal scaling behavior has been analyzed for throughput but not validated at full enterprise scale. The distributed deployment topology is recommended for environments exceeding approximately 10 concurrent agents.

Endpoint monitoring coverage gaps. osquery must be installed on each monitored host. Hosts that are not instrumented — including some cloud-hosted execution environments — are invisible to the endpoint layer. In such environments, network and identity monitoring carry additional weight.

Credential-free Shadow servers. MCP servers operating without any authentication generate no identity-layer signals. For these servers, detection relies entirely on network and endpoint monitoring, which are effective for locally deployed instances but may miss remote credential-free servers that use standard ports.

Advanced adversarial scenarios. The most sophisticated attack class — an adversary who perfectly mimics a legitimate server across all three monitored dimensions simultaneously — can partially evade the current framework. Agent-layer monitoring with PTAS is the natural extension to address this gap, as discussed in Section 11.

10.3 Deployment Guidance

Based on the testbed experience, the author offers three practical recommendations for organizations deploying SMDGF. First, begin with Zeek network monitoring and osquery endpoint monitoring only, in alert-only mode, for a minimum of two weeks before enabling any automated blocking. This establishes confidence in the false positive rate before policies have operational consequences. Second, use the DMR review workflow as an opportunity to engage — not penalize — teams whose Shadow MCP servers are

discovered. Most Shadow MCP servers are Type I deployments created for legitimate productivity reasons, and the review process should lead to proper registration rather than simple blocking. Third, prioritize identity monitoring deployment for teams using third-party MCP services, where the network fingerprint of Shadow and sanctioned servers may be difficult to distinguish.

10.4 Ethical Considerations

The monitoring sources in SMDGF collect data that could indirectly reveal information about the work patterns and tasks of employees who use AI agents. Organizations deploying the framework should implement data minimization — retaining derived risk scores and statistical summaries rather than raw log content wherever possible — apply access controls limiting telemetry access to authorized security personnel, define and publish clear data retention limits, and communicate transparently to staff that the framework monitors AI infrastructure, not individual employee behavior. SMDGF should not be repurposed as an employee productivity monitoring tool.

11. FUTURE WORK

- **Agent-layer monitoring and PTAS.** As discussed in Section 5.4, a fourth monitoring layer that measures the semantic alignment between agent tasks and tool invocations would extend detection to the most sophisticated adversarial scenarios. Future work will implement PTAS, validate it against controlled prompt injection scenarios, and evaluate whether its additional detection coverage justifies the instrumentation overhead in production environments.
- **Federated baseline sharing.** Organizations running the same well-known MCP packages could share anonymized baseline behavioral profiles, enriching each organization's deviation detection without exchanging sensitive telemetry. A federated learning approach to baseline sharing across enterprises is a natural direction.
- **Broader tool integration coverage.** As the enterprise AI agent ecosystem matures, other integration standards — OpenAI Function Calling, Google Vertex AI extensions, and custom API wrappers — present the same Shadow IT governance challenges as MCP. Extending SMDGF's monitoring and registry architecture to cover these additional integration patterns would provide comprehensive coverage of the agentic tool attack surface.
- **Automated remediation workflows.** The current framework detects and alerts. Future work should integrate the DMR with IT service management platforms to initiate structured investigation and remediation workflows automatically, reducing the time from detection to resolution.
- **Large-scale deployment validation.** Validating SMDGF in a real enterprise environment at scale — with hundreds of agents, dozens of MCP servers, and real organizational governance processes — would substantially strengthen the evidence base for the framework's practical deployability.

12. CONCLUSION

This paper introduced Shadow MCP Servers as a distinct and consequential security risk in enterprise AI deployments: a risk that is real, growing rapidly alongside the adoption of LLM agent systems, and not addressed by any existing enterprise security product. The paper formalized the problem through a four-class origin taxonomy, explained why AI agent autonomy makes this risk qualitatively different from

traditional Shadow IT, and presented SMDGF as a practical, deployable detection and governance framework.

SMDGF combines network traffic analysis, endpoint process monitoring, and identity and access telemetry into a unified risk scoring system that continuously discovers MCP infrastructure and drives a detection-driven governance registry. The Dynamic MCP Registry translates detection signals into structured, auditable governance workflows and enforces least-privilege access policies between AI agents and the tools they invoke. All three monitoring sources are built on widely deployed, production-proven infrastructure — Zeek, osquery, and enterprise identity providers — making SMDGF adoptable without requiring significant new infrastructure investment.

Empirical evaluation on a real enterprise-equivalent testbed demonstrates $F1 = 1.00$ with median discovery latency of **36** seconds across all four Shadow MCP threat classes, substantially outperforming three baseline approaches while imposing less than **2%** processing overhead. The framework is released as open source to support community validation and to accelerate adoption in enterprises where AI agent governance is an immediate and growing need.

DATA AVAILABILITY

The experimental dataset supporting the findings of this study, including Zeek network logs, osquery endpoint telemetry, ground truth injection records, and the SMDGF scorer implementation, will be made available by the author upon reasonable request. Requests may be directed to the corresponding author via the journal editorial office.

REFERENCES:

- [1] Anthropic, "Model Context Protocol Specification v1.0," Open Source Project, November 2024. [Online]. Available: <https://modelcontextprotocol.io>
- [2] MCP Community, "MCP Server Registry," GitHub Repository, 2025. [Online]. Available: <https://github.com/modelcontextprotocol/servers>
- [3] A. Behrens, J. Morisse, and C. Meske, "Shadow IT: A Systematic Literature Review," in Proc. 25th European Conference on Information Systems (ECIS), Guimarães, Portugal, 2017.
- [4] S. Silic and A. Back, "Shadow IT — A View from Behind the Curtain," *Computers & Security*, vol. 45, pp. 274–283, Elsevier, 2014.
- [5] S. Rose, O. Borchert, S. Mitchell, and S. Connelly, "Zero Trust Architecture," NIST Special Publication 800-207, National Institute of Standards and Technology, Gaithersburg, MD, 2020.
- [6] Y. Ruan et al., "Identifying the Risks of LM Agents with an LM-Emulated Sandbox," arXiv preprint arXiv:2309.15817, 2023.
- [7] K. Greshake, S. Abdelnabi, S. Mishra, C. Endres, T. Holz, and M. Fritz, "Not What You've Signed Up For: Compromising Real-World LLM-Integrated Applications with Indirect Prompt Injection," in Proc. ACM Workshop on AI and Security (AISec), Copenhagen, 2023.
- [8] National Institute of Standards and Technology, "Artificial Intelligence Risk Management Framework (AI RMF 1.0)," NIST AI 100-1, Gaithersburg, MD, January 2023.
- [9] OWASP Foundation, "OWASP Top 10 for LLM Applications 2025," Version 2.0, 2025. [Online]. Available: <https://owasp.org/www-project-top-10-for-large-language-model-applications>



- [10] V. Paxson, "Bro: A System for Detecting Network Intruders in Real-Time," Computer Networks, vol. 31, no. 23–24, pp. 2435–2463, 1999. (Now: Zeek — <https://zeek.org>)
- [11] Facebook (Meta), "osquery: SQL Powered Operating System Instrumentation," Open Source Project, 2024. [Online]. Available: <https://osquery.io>
- [12] Open Policy Agent, "OPA: Policy-Based Control for Cloud Native Environments," CNCF Graduated Project, 2024. [Online]. Available: <https://www.openpolicyagent.org>
- [13] European Parliament, "Regulation (EU) 2024/1689 — Artificial Intelligence Act," Official Journal of the European Union, June 2024.