

ShardSecure MCP



Secure Context Enrichment for AI Inference

The quality and relevance of an AI model's responses depend heavily on the context available at inference time, including the prompt, session history, and any external data provided. Improving this context directly improves the usefulness and accuracy of the model's output.

ShardSecure MCP (model context protocol) is designed to securely enrich AI context by allowing models to retrieve only the minimum required, authorized data fragments from encrypted storage, leveraging rich metadata stored at the object level to enable precise, efficient searches without needing to scan or retrieve unnecessary portions of the entire dataset — without ever exposing raw data stores to users or AI systems.



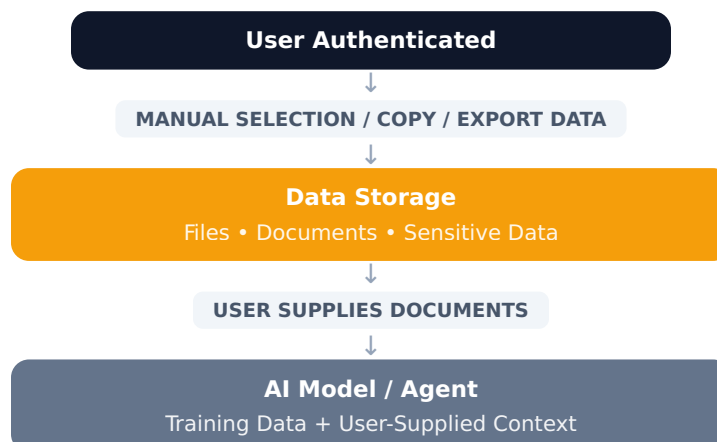
Traditional Approach: Without ShardSecure MCP

When an authenticated user interacts with an AI agent service or chat interface, the model often requires additional context to answer accurately. In the current scenario:

The user must manually locate, select, and provide relevant files or data. The AI relies on a limited combination of:

- Existing training data
- Static or always-on external sources
- User-supplied documents
- Users typically require direct, human-readable access to the data
- Sensitive information must often be copied, exported, or duplicated

Traditional Workflow





Secure Approach: With ShardSecure MCP

Instead of manual data handling, the process becomes seamless and secure:

- The authenticated user provides a natural prompt containing keywords, identifiers, or index references.
- The AI model uses these signals to query the ShardSecure MCP.

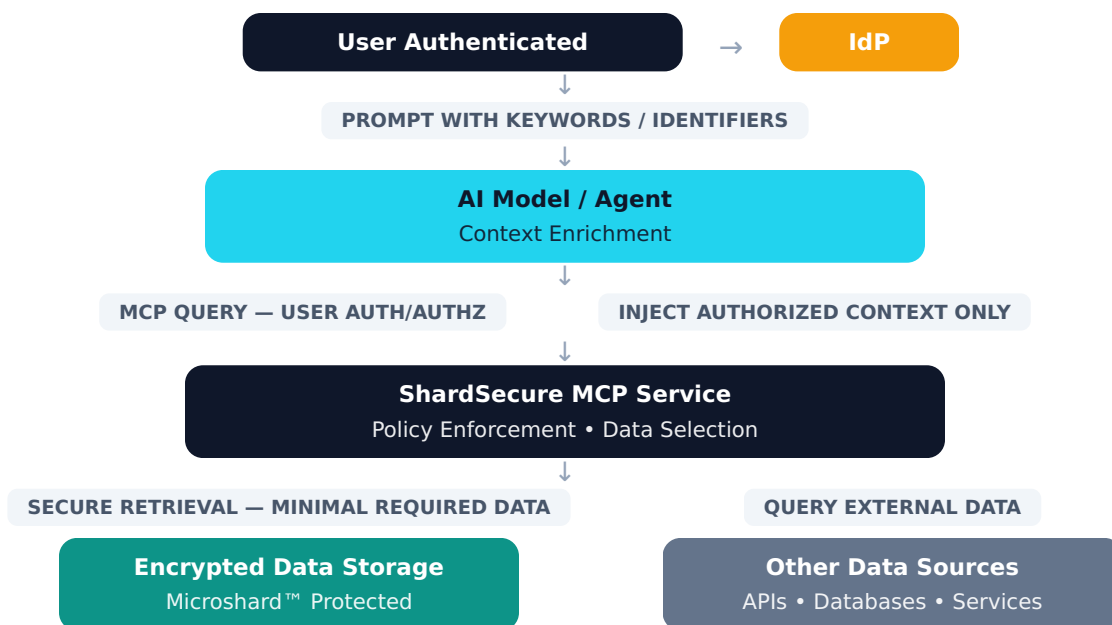
ShardSecure MCP :

- Applies the user's and agent's authentication and authorization policies
- Selects only the precise, authorized data fragments from encrypted storage
- Injects the minimal required information directly into the model's context

The user does not need direct access to the underlying data and may not even be aware of which additional context was retrieved. If more information is needed, the AI can automatically issue additional MCP requests, or request explicit user approval, depending on policy configuration.

Crucially: at no point does the AI model or the user access the raw data store. All data remains encrypted within the ShardSecure platform. Only the specifically authorized data required for the task is shared.

Secure Workflow



Key Benefits

- ✓ **Reduced data exposure** — Sensitive information stays protected and is shared only in minimal, need-to-know fragments. Prevents leakage of PII and PHI.
- ✓ **Stronger access controls** — Granular, policy-driven authorization at the data-fragment level.
- ✓ **Zero trust** — Enforces least-privilege access at the data-fragment level, reducing compliance and audit risk.
- ✓ **Dynamic masking** — Enables granular, per-user data masking.
- ✓ **Less manual data handling** — Eliminates tedious searching, copying, and exporting.
- ✓ **More relevant AI responses** — Higher-quality context leads to significantly better model performance.



Example Use Cases

INSURANCE

1. Fraud Detection and Claims Trend Analysis

Scenario: An insurance company needs to analyze thousands of claims documents, photos, medical reports, and invoices to detect fraud patterns or generate actuarial statistics (e.g., claim frequency by vehicle type, injury category, or region).

Without ShardSecure: Feeding the full unstructured dataset (sensitive personal and financial details) into an AI tool via MCP risks data leakage, regulatory violations (e.g., GDPR, state insurance laws), and high transfer costs. Insurers often avoid AI altogether or use heavily anonymized subsets that reduce accuracy.

With ShardSecure MCP: Claims files are microsharded and stored with rich object-level metadata (e.g., claim type, amount ranges, red-flag indicators, policy categories, timestamps). The AI queries the service for fraud statistics or trends; the mediation layer uses metadata to identify and pull only the minimal relevant data. Raw documents and photos stay protected and fragmented — the model gets aggregated insights or anonymized patterns only. **Benefit:** Faster, more accurate fraud detection and analytics without ever granting the AI, or analysts, direct access to sensitive customer files.

HEALTHCARE

2. Population Statistics from Large X-ray / Medical Imaging Repositories

Scenario: A hospital network or research institute wants to query millions of X-ray, MRI, or CT images to generate aggregate statistics (e.g., prevalence of specific conditions by age group, regional trends, or treatment efficacy metrics) for public health studies or internal reporting.

Without ShardSecure: The AI model, via a standard MCP, would require access to large batches of raw image files or full datasets to compute statistics — exposing sensitive patient data (PHI), massive data transfers, and high compliance risk under HIPAA/GDPR. A breach during the process could expose identifiable images.

With ShardSecure MCP: The AI requests statistics through the service. Using rich object-level metadata (e.g., anonymized tags for condition indicators, scan type, patient demographic buckets, timestamps, or pre-computed features), the system filters relevant microsharded objects first, then retrieves and processes only the necessary tiny shards/fragments for the calculation — never exposing raw images or full patient records. **Result:** Secure, low-bandwidth analysis with full auditability and zero raw data exposure.