

Section 7 — Case Study D (Empirical)

Infrastructure Exposure at AI-Scale Data Gravity: The DeepSeek ClickHouse Database Incident (January 2025)

Drop-in section for the AISIF Practitioner White Paper · Formatted to the Section 7 case-study convention · Prepared July 2026

Unlike Case Studies A–C, which are illustrative composites, Case Study D analyses a publicly disclosed incident. It therefore serves a second purpose identified in Section 8 (Future Work): providing empirical grounding for the framework's analytical claims by mapping a real-world incident to AISIF layers, control domains and maturity-level indicators.

Case Study Card D — presented in the standard Section 7 card format.

Case Study D — Infrastructure Exposure: DeepSeek ClickHouse Database AI Model Provider (empirical, disclosed January 2025)	
Exposure vector	Publicly accessible ClickHouse database on internet-facing subdomains (ports 8123 and 9000), reachable without any authentication. The HTTP interface permitted arbitrary SQL execution from a browser, granting read access — and full database control — over more than one million log-stream entries containing plaintext chat history, API secrets, backend details and operational metadata. Discovered externally by Wiz Research within minutes of a routine posture assessment; no adversarial AI technique was required or used.
AISIF layers implicated	Operations Layer (primary) — unauthenticated, unsegmented, internet-exposed datastore; no attack-surface monitoring; detection occurred by external disclosure, not internal telemetry. Data Layer (primary amplifier) — inference logs treated as operational exhaust rather than classified user data; chat content and secrets logged in plaintext with no minimization or redaction at ingestion. Governance Layer (root cause) no pre-production exposure gate operating at deployment velocity; asset inventory and risk assessment outpaced by viral product adoption. Application Layer (secondary) secrets persisted through the logging pipeline rather than vaulted; exposed API keys created downstream compromise paths. Model Layer — not directly attacked; latent upward risk propagation only (credentialed abuse of inference endpoints enabled by exposed secrets).
Failure mode	Governance-velocity failure manifesting as Operations-layer control absence, with Data-layer classification failure as the amplifier. The proximate cause (missing access control) sits two layers below the framework-relevant causes. No AI-specific attack surface was involved: this was a conventional cloud misconfiguration amplified by AI-scale data gravity, in which the inference traffic itself constituted the sensitive data.
What single-framework analysis missed	The OWASP LLM Top 10 contains no category that squarely captures an unauthenticated ancillary datastore; MITRE ATLAS models adversarial techniques, and none was employed; NIST AI RMF names the governance gap but does not operationalize exposure management. The controls that would have prevented the incident are ISO/IEC 27001 Annex A fundamentals —

	network access restriction, secure configuration, asset management — demonstrating that AI-specific frameworks without the ISO 27001 technical baseline lack operational teeth.
AISIF mitigation	Operations: authentication and network isolation on all AI-adjacent datastores; continuous external attack-surface monitoring; exposure detection integrated into the recursive feedback loop. Data: classify inference logs at the sensitivity of their contents; redact chat content and secrets at log ingestion; minimization by default. Governance: mandatory pre-production exposure review gate; AI infrastructure asset inventory maintained at deployment cadence. Application: vaulted, short-lived secrets; prohibition on credentials transiting the logging pipeline.
Standards activated	ISO/IEC 27001:2022 Annex A (access control, secure configuration, network security, asset management) · NIST AI RMF Govern and Manage functions · NIST AI 600-1 (data privacy and information security actions) · CSA AI Controls Matrix (infrastructure and data-security domains) OWASP LLM02 Sensitive Information Disclosure (partial, downstream).

7.4.1 Incident Overview

In late January 2025, researchers at Wiz conducted an external security-posture assessment of DeepSeek, a Chinese AI model provider whose DeepSeek-R1 reasoning model had achieved viral adoption in the preceding days. Within minutes, the researchers identified two unusual open ports (8123 and 9000) on internet-facing DeepSeek subdomains, leading to a publicly accessible ClickHouse database that required no authentication of any kind. ClickHouse's HTTP interface exposed a /play path permitting arbitrary SQL execution directly from a browser.

A log_stream table within the database contained more than one million entries, including plaintext chat history, API secrets, backend service details and operational metadata, with records dating from early January 2025. Beyond read access, the configuration permitted full control of database operations, creating potential for data modification, deletion and privilege escalation within the DeepSeek environment. Wiz disclosed the exposure responsibly; DeepSeek secured the database promptly. There is no public evidence establishing whether the data was accessed by other parties before disclosure, and no public evidence of the systemic remediation measures taken beyond closure of the exposure itself.

Two properties make this incident analytically significant for AISIF. First, no AI-specific technique was involved: no prompt injection, no model extraction, no adversarial input. Second, the exposed asset was not a primary application database but a telemetry store — yet in an AI system, the telemetry is the user data. The incident is therefore a conventional cloud security failure amplified by AI-scale data gravity, and it stresses precisely the seam between traditional and AI-specific frameworks that AISIF exists to close.

7.4.2 Layer-by-Layer Analysis

Table D-1. AISIF layer implication analysis.

AISIF Layer	Implication	Analysis
Governance	Root cause	Deployment velocity outran the risk-assessment cycle: the database sat on public subdomains, indicating either an incomplete inventory of internet-facing AI infrastructure or an inventory without a pre-production exposure-review requirement. The Governance layer's function of imposing mandatory gates on the layers below was not operating at the cadence of product adoption.
Data	Primary amplifier	The exposed asset was implicitly classified as operational exhaust rather than as user data. Two distinct control failures: data classification did not extend to derived and ancillary stores, so logs inherited none of the protection their contents warranted; and no minimization or redaction operated at log ingestion, so chat content and API secrets were persisted in plaintext at all. Principle for the framework: log pipelines in AI systems are data stores of the same sensitivity class as the primary application database, because inference traffic is the user data.
Model	Not attacked; latent	The model itself was untouched — the analytically important observation. Exposed API secrets and backend details nevertheless created second-order Model- and Application-layer risk: credentialed abuse of inference endpoints and adversary insight into orchestration. AISIF's layered view

		captures this as upward risk propagation from an Operations-layer breach into layers that were never directly attacked.
Application	Secondary	Secrets-management failure: API keys transited and were persisted by the logging pipeline rather than being vaulted, short-lived and redacted. The application stack treated credentials as logable strings.
Operations	Primary failure locus	Every proximate cause resides here: an unauthenticated database on public ports, no network segmentation or ingress restriction, no continuous attack-surface monitoring, and detection by external disclosure rather than internal telemetry. All are pre-AI controls within ISO/IEC 27001 Annex A territory — network security, secure configuration, access control, asset management.

7.4.3 Control Domain Mapping

Table D-2 maps each observed failure to the AISIF control domain and the operative external framework reference, following the Table 3 crosswalk convention. Control-domain identifiers should be aligned to the taxonomy numbering in Section 4 at integration.

Table D-2. Failure-to-control-domain crosswalk.

Observed failure	AISIF layer	Control domain	Framework reference
Unauthenticated internet-facing datastore	Operations	Infrastructure access control	ISO 27001 A.5.15, A.8.20–8.22; CSA AICM infrastructure domain
No attack-surface monitoring; external discovery	Operations	Exposure management and monitoring	ISO 27001 A.8.16; NIST AI RMF Manage; ATLAS reconnaissance techniques
Plaintext chat history and secrets in logs	Data	Data classification and minimization	NIST AI 600-1 data-privacy actions; CSA Data Security within AI Environments; OWASP LLM02 (downstream)
Secrets transiting the logging pipeline	Application	Secrets management	ISO 27001 A.5.17; OWASP LLM guidance on credential handling
No pre-production exposure gate at scale velocity	Governance	AI asset inventory and risk gating	NIST AI RMF Govern; ISO/IEC 42001 management-system clauses
Remediation confirmed; systemic learning unevidenced	Governance / Operations	Recursive feedback loop	AISIF feedback-loop requirement (Section 3); NIST AI RMF Manage

7.4.4 Maturity-Level Indicators

Case Study D provides the strongest available demonstration that AI security maturity is layer-differentiated. The organization operated at a demonstrably high capability level in model development while exhibiting Level 1 (Ad Hoc) indicators at the Operations layer. An aggregate maturity score would

have concealed this asymmetry; a per-layer maturity profile makes it visible and actionable. If the AISIF assessment instrument does not already score each layer independently, this case constitutes the empirical justification for doing so.

Table D-3. Observed maturity indicators by layer.

AISIF Layer	Observed indicator	Indicative level
Governance	No evidence of a pre-production exposure gate operating at deployment cadence; security posture reactive to external disclosure.	Level 1 — Ad Hoc
Data	No classification of ancillary stores; sensitive content persisted in plaintext telemetry.	Level 1 — Ad Hoc
Model	High engineering capability; security-specific maturity at this layer not observable from the incident.	Not assessable
Application	Credentials handled as loggable strings; no vaulting evident in the exposure path.	Level 1–2
Operations	Unauthenticated public datastore; no internal detection; person-independent monitoring absent.	Level 1 — Ad Hoc

Assessment caveat: these indicators are inferred from a single disclosed exposure and its public reporting. They characterize the state of the affected deployment path at the time of the incident, not the organization’s program in its entirety.

7.4.5 Mitigations

Mitigations are stated in AISIF layer order, with the recursive feedback loop as the closing requirement.

- **Governance** — institute a mandatory exposure-review gate for any AI-adjacent infrastructure reaching an internet-facing state, with the gate's cadence bound to deployment cadence rather than to a periodic review calendar; maintain the AI infrastructure inventory as a deployment artefact, not an audit artefact.
- **Data** — classify inference logs, traces and telemetry at the sensitivity of their contents; enforce redaction of chat content, personal data and secrets at the point of log ingestion; adopt minimization by default for all derived stores.
- **Application** — vault all credentials with short-lived issuance; prohibit secrets from transiting prompt, log or trace pipelines; rotate any credential class that has ever appeared in telemetry.
- **Operations** — require authentication and network isolation on every datastore in the AI serving path, including analytics and observability stores; operate continuous external attack-surface monitoring so that exposure is detected internally before it is detected externally.
- **Feedback loop** — route the incident finding into the governance cycle with authority to trigger control updates: audit all peer datastores, amend the logging standard, and evidence the systemic change. Closure of the exposure is a response; only the systemic change constitutes learning.

7.4.6 Analytical Significance for the Framework

Case Study D completes the analytical range of Section 7. Case Studies A and B demonstrate failures that traditional frameworks miss and AI-specific frameworks catch. Case Study D demonstrates the inverse: a failure that AI-specific frameworks miss and traditional frameworks catch. The OWASP LLM Top 10 has no category that squarely fits an unauthenticated ancillary datastore; ATLAS models adversarial technique, and none was used; the NIST AI RMF names the governance gap but does not operationalize exposure management. The preventing controls are ISO/IEC 27001 fundamentals.

The incident therefore furnishes direct empirical support for the paper's central claim: an organization implementing AI-specific frameworks without the ISO 27001 technical control baseline has a governance cycle without operational teeth, and an organization with only the traditional baseline lacks the AI-specific lens that recognizes inference telemetry as crown-jewel data. Integration, not adoption, is the requirement — and the disclosed incident record to date suggests that AI security incidents are disproportionately conventional infrastructure failures amplified by AI-scale data gravity.

Sources and Verification Note

Primary source: Wiz Research disclosure, 'Wiz Research Uncovers Exposed DeepSeek Database Leaking Sensitive Information, Including Chat History' (wiz.io/blog, 29 January 2025). Corroborating reporting: TechTarget, TechRepublic, and contemporaneous coverage (January–February 2025). Factual claims in this case study are limited to the published disclosure; where the public record is silent — prior third-party access, and the scope of systemic remediation — the case study states the absence of evidence rather than an inference. Verify quotations and technical particulars against the Wiz publication before print.