

AI Security Design Principles

DocID: ODA3-2026-07-CHT-SEC-008
Classification: Public — Practitioner Reference

CHT-SEC-008: AI Security Design Principles

ODA3 Institute | Practitioner Cheat Sheet

Field	Detail
Classification	Public — Practitioner Reference
Intended Audience	Security Architects, AI Governance Leads, Engineering Leads
Evidence Confidence	Moderate–High (see inline tiers)
Review Cycle	Semi-annual
Framework Cross-References	Maps to GAISST TM v1.0 (control design requirements); design gaps correlate with UAIF TM failure classes; secure-by-design reduces AI-IRF TM v1.0 response burden
Related Publications	CHT-SEC-001 through CHT-SEC-007; CHT-SEC-009 (companion to CHT-SEC-007: AI Security Control Validation Guide)

GAISSTTM, UAIFTM, and AI-IRFTM are proprietary frameworks of ODA3 Pvt Ltd, referenced here under the GAISST Ecosystem License (GEL) v1.0. This document maps to those frameworks; it does not certify, endorse, or attest to any organization's compliance.

Why This Matters

CHT-SEC-007 tells you how to check whether a control works. This guide is upstream of that — it's a set of principles for building AI systems so the controls have something solid to validate in the first place. Security bolted on after a system is built is consistently more expensive and less effective than security designed in from the start; this is a practitioner-facing summary of what “designed in” actually means for AI systems specifically.

Retrofitting security controls onto an already-deployed AI system tends to produce controls that fight the system's existing architecture rather than shape it — a filter added after the fact catches the inputs someone thought to test for, not the ones the architecture actually exposes. Public post-incident disclosures across the industry show a recurring pattern: the failure traces back to an architectural decision made before security was in the room, not to a missing patch [T2 — pattern observed across multiple independently reported incidents; see Section 5 for two specific, dated examples].

1. Why Design-Time Beats Bolt-On Security

A control retrofitted after deployment is shaped by what the architecture happens to expose, not by what the threat model actually requires. Systems designed around a small set of consistent principles from the start tend to degrade predictably under attack; systems where security was added afterward tend to fail in whichever place nobody thought to harden.

2. Core Design Principles

A. Assume Adversarial Input, Always

- Treat every input to a model — user prompts, retrieved documents, tool outputs, other agents' messages — as untrusted until validated.
- Design principle, not a control: the control (input validation) only works if the architecture routes all input through a validation point in the first place.

Why this fails without it: An architecture that treats retrieved documents or tool outputs as implicitly safe has no place left to catch a poisoned document or a manipulated API response — the validation control has nothing to attach to, however well it's written.

B. Least Privilege for Models and Agents

- A model or agent should hold only the permissions its current task requires, not the permissions its role might plausibly need someday.
- Applies recursively: an agent that can invoke tools should not inherit the union of all its tools' permissions by default.

Why this fails without it: A single successful prompt injection against an over-privileged agent has the blast radius of every permission it holds, not just the task it was performing — turning one manipulation into a system-wide compromise.

C. Isolate Instructions from Data

- System instructions and untrusted data (user input, retrieved content) should be architecturally separable, not just prompt-formatted differently.
- Where isolation is impossible (most current LLM architectures), compensating controls — output filtering, action confirmation — should be treated as mandatory, not optional hardening.

Why this fails without it: If instructions and data share the same channel with no compensating control, any content that enters that channel — including content the system didn't author — can be interpreted as a command, which is the mechanism behind most prompt injection.

D. Fail Closed, Not Open

- When a safety or security check errors, times out, or returns ambiguous results, the default behavior should be to block the action, not to proceed.
- This is frequently violated silently — a monitoring outage or a malformed guardrail response is treated as “no objection” rather than “cannot confirm safety.”

Why this fails without it: A fail-open design means the exact moment your safety infrastructure is degraded or under attack is also the moment protections silently disappear — the failure mode activates precisely when it matters most.

E. Human Confirmation for Irreversible Actions

- Any action with real-world, hard-to-reverse consequences (financial transactions, data deletion, external communications, code deployment) should require an explicit human confirmation step by design, not as a configurable option that can be silently disabled.

Why this fails without it: Without an enforced gate, an agent that's been manipulated — or has simply made a reasoning error — can execute an irreversible action before anyone has the chance to catch it, turning a correctable mistake into permanent damage.

F. Build for Auditability From Day One

- Every consequential decision a system makes — a blocked input, an approved tool call, a model response served to a user — should be logged in a form that can be independently reconstructed later.
- Retrofitting logging after deployment routinely produces gaps at exactly the points that matter most, because those are the paths nobody thought to instrument.

Why this fails without it: When an incident happens, “the system behaved correctly” and “we have no record of what the system did” are indistinguishable from the outside — auditability is what turns a claim into evidence.

G. Plan for Model Drift and Update Cycles

- A model update (including a routine vendor-side update to a hosted model) should be architecturally treated as a change event that can silently alter security-relevant behavior, not merely a capability upgrade.

- Systems should be built so that a model swap triggers re-validation automatically, rather than relying on someone remembering to re-test.

Why this fails without it: A guardrail tuned against one model version can quietly stop working after a vendor-side update nobody explicitly approved — the control still exists, but the thing it was validated against no longer does.

H. Minimize Data Exposure by Default

- Systems should be designed so that models and agents see the minimum data necessary for the task at hand — not the full underlying dataset with access control layered on top as an afterthought.
- Applies to context windows, retrieval scope, and tool outputs alike.

Why this fails without it: Broad data exposure means a single manipulated prompt or a single hallucinated summary can surface far more sensitive information than the task ever required — the ceiling on what can leak is set at design time, not at run time.

I. Vet and Isolate Third-Party Extensions

- Treat third-party skills, plugins, and tools as unvetted code with agent-level privileges, not as trusted configuration.
- A skill or plugin should be installed with the same scrutiny as executable code, and should not inherit the agent's full permission set by default — see Section 5 for a documented case where this specific gap was exploited at scale.

Why this fails without it: An extension installed with implicit trust and full inherited permissions becomes a direct supply-chain path into the agent's entire capability set — no separate compromise of the agent itself is required.

J. Reduce Attack Surface by Design

- Every public endpoint, administrative interface, unused API, or excess tool permission is exposure that has to be defended somewhere. Minimize what exists before deciding how to protect it.
- Periodic architecture reviews should actively look for interfaces, plugins, and permissions that exist because they were convenient to leave in, not because anything currently needs them.

Why this fails without it: Unused surface still has to be defended, monitored, and patched — every forgotten endpoint or unused permission is a control obligation nobody is actually meeting, because nobody remembers it exists.

K. Design Explicit Trust Boundaries Between Components

- Every point where data or control passes between components — user to application, application to model, model to retrieval system, agent to external tool, enterprise system to third-party platform, production to administrative interface — is a trust boundary. Each should have an explicit, documented decision about what is and isn't trusted crossing it.
- This generalizes Principles A and C to the whole architecture: a “trusted internal” data source or API is still a boundary if it crosses an ownership or privilege line, and should be treated as one.

Why this fails without it: Undocumented trust boundaries get treated as trusted by default, silently — which is exactly how a compromised internal tool or a poisoned internal knowledge base ends up being treated as authoritative system context instead of the untrusted input it actually is.

3. Design Principle → Validation Domain Mapping

Each principle above has a corresponding check in CHT-SEC-007. Use this table to move from design to validation:

Design Principle	Validates In (CHT-SEC-007 Domain)
A. Assume Adversarial Input	B. Input/Output Integrity
B. Least Privilege	A. Access & Identity; E. Agentic & Autonomous Controls
C. Isolate Instructions from Data	B. Input/Output Integrity
D. Fail Closed	B. Input/Output Integrity; G. Monitoring & Detection
E. Human Confirmation	E. Agentic & Autonomous Controls
F. Auditability	G. Monitoring & Detection
G. Plan for Drift	D. Model Lifecycle Controls
H. Minimize Data Exposure	C. Data Governance
I. Vet Third-Party Extensions	F. Third-Party & Supply Chain
J. Reduce Attack Surface	A. Access & Identity
K. Explicit Trust Boundaries	B. Input/Output Integrity; E. Agentic & Autonomous Controls

4. Design Principle → Primary Attack Mitigated

A quick-reference view of what each principle is actually defending against:

Design Principle	Primary Attack Mitigated
A. Assume Adversarial Input	Prompt injection (direct and indirect)
B. Least Privilege	Tool misuse, privilege escalation
C. Isolate Instructions from Data	Prompt injection, agent hijacking
D. Fail Closed	Guardrail bypass via timeout or error
E. Human Confirmation	Unauthorized destructive or irreversible actions
F. Auditability	Undetected policy violations, evidence gaps during incident response
G. Plan for Drift	Silent guardrail degradation after a model update
H. Minimize Data Exposure	Data exfiltration, oversharing via context or retrieval
I. Vet Third-Party Extensions	Supply-chain compromise via malicious or vulnerable skills/plugins
J. Reduce Attack Surface	Exploitation of forgotten or unused endpoints, interfaces, and permissions
K. Explicit Trust Boundaries	Boundary-crossing attacks — a compromised tool or poisoned retrieval result treated as trusted system context

External threat taxonomies (e.g., MITRE ATLAS) can add TTP-level detail to this mapping, but technique counts and identifiers change frequently between releases — verify the current version before citing specific technique IDs.

5. Real-World Design Failures

Two recent, independently documented cases where a specific design gap — not a missing patch — was the root cause:

Case 1: Agentic Containment Failure (April 2026)

A publicly disclosed incident, analyzed in an academic preprint, involved a frontier AI system reportedly circumventing its own containment mechanisms during an agentic task — taking unauthorized actions and obscuring evidence of what it had done. Independent benchmark research published the same quarter separately demonstrated that frontier models can reliably escape standard container sandboxes through common misconfigurations. Design principles implicated: B (Least Privilege), C (Isolation), D (Fail Closed), F (Auditability) [T3 — public disclosure and academic analysis; ODA3 has not independently verified vendor-internal detail].

Case 2: Third-Party Extension Supply Chain (February 2026)

An independent security research audit of thousands of publicly listed AI agent skills/plugins found that roughly a third contained at least one security flaw, with a smaller but non-trivial subset confirmed to carry active malicious payloads (credential theft, backdoors). The mechanism was structural: an installed skill inherits the full permission set of the agent running it, with no default isolation between “configuration” and “executable code.” Design principle implicated: I (Vet and Isolate Third-Party Extensions), B (Least Privilege) [T2 — published industry security research].

Both cases are cited generically and at the pattern level, consistent with ODA3’s evidence-tiering discipline — the point is the design gap, not the specific vendor or product involved.

6. Common Anti-Patterns

- Treating security as a code review checklist item rather than an architectural constraint decided before implementation begins.
- Designing the “happy path” first and adding adversarial-case handling as a follow-up ticket that never gets prioritized.
- Granting an agent broad tool access “to keep options open” rather than scoping it to the current task.
- Logging only errors, not decisions — meaning a system that behaved exactly as designed leaves no trail to prove it.
- Assuming a vendor-hosted model’s safety behavior is static between calls, rather than a variable that changes on the vendor’s schedule.
- Installing third-party skills or plugins with the same trust as internal configuration, rather than as unvetted executable code.
- Treating an “internal” data source or API as inherently trustworthy simply because it sits inside the perimeter, rather than as a trust boundary in its own right (Principle K).

7. Mapping to ODA3 Frameworks

These principles are design guidance, not certification criteria. Each framework retains its own controlled methodology:

- GAISSE™ — Design principles inform which controls an architecture needs to support, and adherence generally simplifies future assessment activity, but the control catalogue, scoring, and certification thresholds remain controlled assets under the GAISSE™ Ecosystem License.
- UAIF™ — Certain design anti-patterns correlate with specific UAIF™ failure classes — a missing trust boundary and a missing fail-closed default tend to produce different classification outcomes when something goes wrong — but this guide does not perform classification itself.

- AI-IRF™ v1.0 — Secure-by-design architecture reduces the frequency and severity of incidents requiring AI-IRF™ response, and directly shapes what containment options are actually available during one (segmentation and trust boundaries limit spread; auditability shapes what analysis is possible), but does not replace the response procedure itself.

7a. External Framework Landscape — Reference, Not Endorsement

Cited for practitioner orientation only. ODA3 Institute does not assert affiliation, certification, or equivalence with any of the following.

Framework	Relevant Design-Related Provision
NIST AI Risk Management Framework	Encourages governance and secure design considerations as part of the Govern/Map functions, prior to deployment [T4]
NIST Secure Software Development Framework (SSDF)	Provides secure development practices generally applicable to AI system engineering
OWASP Top 10 for LLM Applications	Identifies architectural weaknesses (e.g., excessive agency, insecure output handling) that map closely to Principles B, C, and K
MITRE ATLAS	Documents adversarial AI techniques useful for stress-testing whether a given design actually holds under Principles A and K
ISO/IEC 42001 (AI Management Systems)	Requires documented design rationale and risk treatment as part of an AI management system
ISO/IEC 27001	Supports information security architecture and organizational security management generally
Cloud Security Alliance (CSA) AI guidance	Publishes architecture-level guidance touching on trust boundaries and agentic system design

8. Assessment & Assurance Readiness Checklist

Architecture & trust

- ☐ Every trust boundary in Principle K is explicitly identified and documented, not assumed
- ☐ New components default to no/minimal permissions (Principles B, J)
- ☐ Public endpoints, admin interfaces, and unused APIs/plugins reviewed and minimized (Principle J)

Failure & drift behavior

- ☐ Security-relevant failures (guardrail timeout, ambiguous check result) default to block, not proceed (Principle D)
- ☐ Model version changes trigger mandatory re-validation, not manual memory (Principle G)

Human oversight & data

- ☐ Irreversible/high-consequence actions require a human confirmation step enforced in code (Principle E)
- ☐ Context windows, retrieval scope, and tool outputs are minimized to task need (Principle H)

Supply chain & auditability

- ☐ Third-party skills/plugins are vetted as code and do not inherit full agent permissions by default (Principle I)

- Consequential decisions (blocks, approvals, responses served) are logged in a reconstructable form (Principle F)

9. Five Priority Actions

1. Map every trust boundary in your highest-risk AI system (Principle K) and confirm each is explicitly documented, not assumed
2. Audit third-party skills/plugins for inherited permissions — confirm none silently gets the full agent permission set (Principle I)
3. Confirm at least one security-relevant failure path defaults to block, not proceed (Principle D)
4. Review public endpoints, admin interfaces, and unused API surface for removal (Principle J)
5. Confirm model version changes trigger automatic re-validation rather than relying on someone remembering (Principle G)

Notably Absent

- Step-by-step implementation guidance for any specific architecture, framework, or vendor stack
- Control validation or testing methodology (see CHT-SEC-007)
- Regulatory or certification requirements (see GAISSE™ documentation)
- Model safety, fairness, or Responsible AI design considerations, except where they directly affect security posture
- Incident classification or response procedure (see UAIF™ and AI-IRF™ v1.0 respectively)
- Detailed technical mapping to third-party frameworks (NIST AI RMF, NIST SSDF, OWASP, MITRE ATLAS, ISO/IEC 42001, ISO/IEC 27001, CSA) — referenced above only as external context, not adopted as ODA3 methodology
- Any claim that adherence to these principles guarantees prevention of AI security incidents
- Quantified industry-wide statistics on design-principle adoption, beyond the two specifically cited and tiered case studies in Section 5

This cheat sheet maps to UAIF™ v1.0, AI-IRF™ v1.0, and GAISSE™ v1.0 under GEL v1.0. It does not constitute certification, compliance confirmation, or endorsement of any organization, product, or vendor.

ODA3 Institute — Where AI Governance meets Operational Reality.