

ODA3 INSTITUTE

ADOPTION GUIDE

The AI Assurance Imperative

A Board Member's Guide to GAISSEF™

Doc ID: ODA3-2026-08-WHP-COM-008

Classification: Public / Freely Distributable

© ODA3 Pvt Ltd. Published by ODA3 Institute.

Executive Summary

As organizations deploy AI systems, boards may need to determine whether AI-related risks are being governed and whether management can demonstrate that governance with evidence. This guide gives board members a small set of questions to ask about AI governance — grounded in what GAISSE™, ODA3 Institute's AI safety and security framework, is designed to help organizations demonstrate. It assumes no technical background and does not require reading any other GAISSE™ material first.

METHODOLOGY NOTE

This guide is introductory and conceptual. It contains no client outcome data, no proprietary telemetry, and no incident statistics. It does not constitute legal, financial, or governance advice, and does not establish conformance, certification, or regulatory compliance for any organization.

1. The Question Boards Are Now Being Asked

For board oversight, a practical question is: how does management know that AI systems are being operated within defined safety and security controls, and what evidence supports that conclusion?

A policy can describe intended governance. Assurance requires evidence showing whether relevant controls have actually been implemented and are operating as represented. The distance between a governance assertion and the evidence supporting it is an assurance problem.

2. What GAISSE™ Is, in Board Terms

GAISSE™ (the Global AI Safety and Security Framework) is a control framework published by ODA3 Institute: a structured way for an organization to document what safety and security controls exist for its AI systems, and what evidence supports each one. It is not a law, and adopting it does not itself satisfy any specific regulation — but it gives management a common structure for answering oversight questions with something more concrete than a policy statement.

GAISSE™ is one of four frameworks ODA3 Institute publishes, covering AI safety and security governance, incident classification, incident response, and — where relevant — the security of AI systems capable of physical action. This guide focuses on GAISSE™ because its governance, control, and evidence architecture is directly relevant to board-level assurance questions.

3. Questions a Board Can Ask Management

- Which AI systems in our organization are currently in scope of any structured safety or security review, and which are not?
- What evidence — beyond policy statements alone — supports the claim that those controls are implemented and operating as represented?
- Who is accountable for AI-specific safety and security controls, and how does that differ from existing IT or cybersecurity accountability?
- If potentially harmful or anomalous AI behaviour were detected today, is there a defined process for determining whether it constitutes an incident, preserving relevant evidence, classifying it where appropriate, and initiating proportionate response?
- Do any AI systems produce physical effects or control physical equipment, and if so, are the additional physical-AI risks and controls explicitly addressed within the organization's assurance scope?

4. A Worked Scenario (Hypothetical)

Consider a hypothetical company whose board asks whether a newly deployed AI claims-handling assistant is "secure." Management's initial answer is a policy reference and an assurance from the software vendor. Applying GAISSF™-style evidence discipline, the board's follow-up questions would instead surface: a documented boundary on what data the assistant can access, a record of human review before its outputs affect a customer, and a defined escalation path if it behaves incorrectly. The board would then ask whether these controls are actually implemented, what evidence demonstrates their operation, and what remains unknown or unverified. The existence of an artifact alone does not establish control effectiveness or system safety.

This scenario is illustrative only. It does not describe an ODA3 client, engagement, or outcome.

KEY PRINCIPLE

Assurance is demonstrated, not asserted. For a board, the practical implication is simple: ask what evidence supports the policy and whether that evidence demonstrates implementation in practice.

BOARD OVERSIGHT PRINCIPLE

Management should be able to explain what is known, what evidence supports it, what remains uncertain, who owns unresolved risk, and what action is being taken. Board oversight does not require directors to become AI security assessors.

5. What This Guide Does Not Do

This guide does not tell a board whether its organization's AI systems are safe, does not establish a legal standard of care, and does not substitute for advice from qualified counsel, risk, or security professionals. It also does not assume the organization has any GAISSF™ work under way — the questions above are useful with or without formal adoption.

NOTABLY ABSENT

This guide does not establish conformance, certification, regulatory compliance, or a legal duty of care. It is not a substitute for legal, risk, or security advice. ODA3 Institute's assessment and certification infrastructure is under development.

Next Reading

- ODA3 Institute in 15 Minutes
- The 10-Minute Introduction to GAISSF™
- Which GAISSF Ecosystem Framework Do I Need?

GAISSF™, UAIF™, AI-IRF™, and PAI-SF™ are frameworks of ODA3 Institute, referenced under the GAISSF Ecosystem License (GEL) v1.0.