

ODA3 INSTITUTE

EXECUTIVE BRIEF

PAI-SF™ v1.0 — Physical AI Security Framework

Executive Brief

ODA3-2026-07-EXB-PAISF-001 | Open / Public — Publication Candidate

© ODA3 Pvt Ltd

Document Control

Doc ID	ODA3-2026-07-EXB-PAISF-001
Doc Type	Executive Brief (EXB)
Framework Tag	PAI-SF™ v1.0
Audience	Boards, Risk Committees, CEOs, CFOs, General Counsel (primary); CISOs and AI Governance Leads (secondary)
Companion Document	ODA3-2026-07-TCR-PAISF-001 (Technical + Compliance Report)
Governing Licence	GAISSF Ecosystem Licence (GEL v1.0)
Access Classification	Open / Public — no Controlled-asset content in this document
Publication State	Scheduled for Publication — 03 August 2026
Publish Date / Review Cycle	Scheduled for publication 03 August 2026; review cycle to be set upon publication approval
Classification	Public-Facing — Scheduled for Publication (03 August 2026)

PAI-SF™ v1.0 — Physical AI Security Framework

Executive Brief

Methodology Note

This brief summarizes ODA3 Institute's Physical AI Security Framework (PAI-SF™ v1.0) for board and executive audiences. Findings are based on analysis of public disclosures, regulatory filings, and controlled simulation — not proprietary incident data, which ODA3, founded in March 2026, does not yet hold. Where evidence is limited, this brief says so directly rather than implying certainty. The companion Technical + Compliance Report (ODA3-2026-07-TCR-PAISF-001) contains the 41-control register summary and inline evidence sourcing for readers who need it.

The Business Case in One Paragraph

Physical AI — robots, autonomous vehicles, drones, and AI-driven industrial equipment — introduces a form of risk that ordinary IT security and AI governance programs are not built to catch: an AI error here doesn't just leak data, it can cause physical injury, equipment damage, or operational shutdown. PAI-SF™ gives your organization a structured way to answer the question a board, insurer, or regulator will eventually ask: **on what basis do we trust this system not to cause harm?**

Why This Matters Now, in Governance Terms

- **Liability exposure is different in kind, not degree.** A compromised chatbot produces bad text. A compromised physical AI system can produce an uncontrolled machine. Board risk oversight and D&O considerations should treat these as distinct risk categories, not the same category at higher severity.
- **No existing framework was purpose-built for this gap.** Established AI governance standards address model and data risk. Established engineering safety standards address hardware fault risk. Neither addresses what happens when an AI model's own decision — not a hardware fault — causes unsafe physical motion. PAI-SF™ exists specifically to fill that

gap. (The companion Technical Report names the specific standards for readers who need them.)

- **Regulatory and insurance expectations are moving toward this ground, not away from it.** Even where no specific physical-AI regulation yet exists, the same evidentiary discipline — documented risk assessment, tested overrides, incident-ready evidence — is what regulators and insurers already expect in adjacent domains (industrial safety, automotive, medical devices).

What PAI-SF™ Actually Covers

PAI-SF™ organizes physical AI risk into 12 areas of accountability, spanning: sensor and perception trust, physical command safety, hardware integrity, autonomy authorization limits, safe-shutdown behavior, human override capability, monitoring and evidence collection, incident response, supply chain assurance, pre-deployment testing, alignment with existing safety standards, and the physical operating environment itself.

Each area carries specific, auditable expectations — not aspirational principles. The companion Technical Report summarizes all 41 specific controls; this brief focuses on what a governance body needs to know.

Governance and Accountability — What Your Organization Needs in Place

PAI-SF™ expects clear, named accountability across these roles, not a single "AI owner" absorbing everything:

Role	What They're Accountable For
Accountable Owner	Overall risk acceptance and deployment authorization
Safety Authority	Safety requirements and stop-operation authority
Security Authority	Security architecture and incident declaration
AI Governance Authority	Model deployment and update approval
Executive Oversight	Risk appetite, resourcing, and escalation

If your organization deploys or is considering deploying physical AI systems, the first governance question is simple: **are these roles named, or is "someone" assumed to be handling it?**

Financial and Risk Framing

PAI-SF™ does not provide a quantified cost model in this version — that would overstate current evidence. What it does provide is a structure for making informed, defensible risk-acceptance decisions:

- **Cost of inaction is asymmetric.** The cost of implementing structured physical AI assurance is bounded and plannable. The cost of an unaddressed physical AI incident — injury, equipment loss, regulatory scrutiny, reputational damage — is not, and is precisely the kind of tail risk boards are expected to have visibility into.
- **Assurance maturity should be assessed per capability area, not as a single organization-wide score.** A company may be well-prepared on sensor integrity but underprepared on incident evidence preservation, or vice versa. PAI-SF™'s maturity model (still in development) is deliberately structured this way rather than collapsing into one misleading overall grade.
- **No conformance or certification currently exists to "achieve."** ODA3 Institute does not currently hold or offer certification authority for PAI-SF™. Any future conformance pathway is

proposed, not active. Organizations should not represent current alignment with PAI-SF™ as certification, compliance, or regulatory approval.

What This Framework Does Not Do

To be direct about limits: PAI-SF™ is a security assurance framework, not a safety framework. It helps identify where a security failure — a compromised sensor, a hijacked command, a tampered model — could create or worsen a safety consequence. It does not replace your organization's safety engineering, hazard analysis, or existing sector-specific safety certification obligations (aviation, automotive, medical device, industrial). That distinction is intentional, not a gap to be filled later: PAI-SF™ is not claiming safety-engineering authority it doesn't have, and your organization's existing safety obligations remain exactly where they are today. PAI-SF™ also does not eliminate physical risk, does not guarantee safety across all operating conditions, and is not endorsed by any regulator or standards body. It is a structured assurance approach, not a substitute for the safety engineering and regulatory compliance your organization already owes in its sector.

The evidence base for physical AI security specifically — as distinct from general robotics or industrial safety — is still developing industry-wide. This framework will be revised as real-world experience accumulates; it should be treated as a serious starting structure, not a finished, field-proven standard.

Recommended Next Step for Your Organization

If your organization operates or is evaluating physical AI systems, the practical next step is a gap assessment against the 12 accountability areas above — not a compliance exercise, but an honest internal check of where accountability, testing, and evidence collection already exist and where they don't. The companion Technical Report provides the detailed control language needed to run that assessment.