

ODA3 INSTITUTE

TCR — TECHNICAL & COMPLIANCE REPORT

UAIF™ v1.0 Technical Specification

Final Publication v1.0

Document ID: ODA3-2026-06-TCR-STD-002

Classification: Standards Development – Technical Reference

Licensed under GAISSE Ecosystem Licence (GEL v1.0)

UNIFIED AI INCIDENT FRAMEWORK

UAIF™ v1.0

Technical Specification — Final Publication v1.0

TCR-STD-002 | Q2 2026 | Normative errata integrated 14 August 2026 | Licensed under GAISSF Ecosystem Licence (GEL v1.0)

Published by ODA3 Institute | Applied Research & Advisory | AI Security | Standards Development | Certification

Status and Finality Notice

UAIF™ v1.0 is a Final Publication. Final status means the v1.0 normative architecture, controlled vocabulary, machine-readable conformance model, and reference calculation semantics are frozen under ODA3 Institute change control. It does not mean that every sector calibration, legal interpretation, integration pattern, or empirical reliability claim has been validated in production.

Sector-specific calibration, independent inter-rater studies, and broad SIEM/GRC interoperability testing remain validation work. Regulatory trigger logic provides technical routing assistance only. Legal counsel SHALL determine definitive compliance obligations, including EU AI Act Article 73 serious-incident reporting and any GDPR, NIS2, DORA, or sectoral obligations.

1. Executive Summary

UAIF™ provides a machine-readable incident interchange, classification, severity scoring, and conformance architecture for AI incidents and vulnerabilities. It is designed to reduce crosswalk friction across security operations, governance, regulatory routing, and assurance workflows.

UAIF™ is	UAIF™ is not
An incident interchange and classification framework for AI systems	A universal AI governance or ethics framework
A structured severity scoring methodology with documented limitations	A safety certification scheme or legal risk assessment methodology
A conformance-testing architecture for testable adoption	A legal determination of regulatory applicability
An interoperability layer for advisory cross-framework mapping	An empirically validated scoring system across all sectors

1.1A Notably Absent

- No empirical claim that default harm weights are sector-validated.
- No certification, audit authority, or licence to claim UAIF Certified status.
- No legal determination of EU AI Act Article 73, GDPR, NIS2, DORA, or sectoral reporting obligations.
- No universal SIEM/GRC integration warranty across target platforms.
- No proprietary sector calibration packages or confidential pilot datasets.
- No claim that Final Publication status establishes empirical deployment maturity or regulator recognition.

1.2 Primary Outputs

Output Field	Type	Purpose
severity_analytical_score	Float 0.0-10.0	Primary analytical metric; pre-impact, pre-presentation differentiation.
severity_presentation_score	Float 0.0-10.0	Operational presentation metric used for triage and regulatory proxy routing.
severity_level	Integer 1-5	Operational triage tier mapped from presentation score.
severity_rationale	JSON object	Sole canonical, machine-auditable explanation and audit trail for every calculation component.
regulatory_trigger	Boolean	Technical routing flag; not legal determination.
trigger_confidence	LOW MEDIUM HIGH	Calibration-aware confidence band.
human_review_recommended	Boolean	Human review flag for uncertain or low-confidence triggers.

2. UAIF 7-Layer Architecture

Layer	Name	Purpose	Representative Fields
L0	Record Identity & Workflow	Record identity, classification, conformance-profile declaration, and lifecycle/workflow state	incident_uuid, record_type, profile, extensions, reporting_timestamp, reporter_confidence_level, remediation_status, deduplication_hash
L1	Causal Separation Chain	Isolate origin from manifestation	root_cause_primary_domain, root_cause_specific_type, root_cause_confidence, exploit_path_description
L2	Severity Scoring	Impact-dominant severity with transparent limits	severity_analytical_score, severity_presentation_score, severity_level, severity_rationale
L3	Extended Harm Characterisation	Immediate, chronic, and realized-harm characterization	harm_acute, harm_chronic_proxy, realized_harm_categories, cumulative_bias_index
L4	Incident–Vulnerability Relationship	Links classified incidents to associated vulnerabilities; captures exposure context	linked_vulnerability_id, exposure_duration_hours
L5	Generative, Agentic & Control Plane	RAG leakage, hallucination, agent escalation, MCP, OAuth, policy enforcement	ai_system_type, non_adversarial_failure_type, agent_escalation_chain, mcp_endpoint_ids, oauth_pivot_chain, policy_enforcement_point
L6	Regulatory & Cross-Sector Metadata	Jurisdictional routing and sector classification	sector_code, jurisdiction_country_code, regulatory_obligations, regulatory_trigger

Architectural decision: prior L7 AI Control Plane Extensions are merged into L5 for API simplicity. TCR-STD-003 SHALL follow this 7-layer architecture. L7 SHALL NOT be treated as a separate conformance layer in UAIF v1.0.

2.1 Record Type Semantics

record_type SHALL take one of two normative values:

- Incident — a realized event involving an AI system that has produced, or is reasonably assessed to have produced, an adverse operational, security, safety, privacy, societal, or other relevant effect.
- Vulnerability — a weakness, exposure, or susceptible condition that could contribute to an adverse event but is not itself evidence that such an event has occurred.

Unclassified signals, alerts, telemetry items, or investigative leads are not a UAIF record_type. They remain in the originating detection or case-management workflow until sufficient information exists to classify the record as Incident or Vulnerability. Observation is not a UAIF record_type value.

2.2 Profile and Extension Declaration

Every UAIF record SHALL declare profile, identifying its base conformance profile. The normative base values are Core, Enterprise, and Regulatory, where Enterprise extends Core with L3/L4 fields and Regulatory extends Enterprise with extended L6 obligations.

A record MAY additionally declare one or more extensions. UAIF v1.0 defines two extension identifiers: SOC_SIEM (approved ATLAS/security-tooling mapping fields) and AI_Security (selected L5 generative/agentic/control-plane fields). Extensions are additive and SHALL NOT reduce or override the requirements of the selected base profile. Any base profile MAY declare either, both, or neither extension.

3. Adoption Profiles and Conformance

Type	Name	Layer coverage / mandatory fields	Prohibited / gated fields
Base profile	Core	L0, L1, L2, L6 baseline; for record_type = Incident, the L2 severity bundle only	L3 and L4 fields in primary payload

Type	Name	Layer coverage / mandatory fields	Prohibited / gated fields
Base profile	Enterprise	Core + L3, L4; for record_type = Incident, L2 bundle + realized_harm_categories	None beyond undeclared extension fields
Base profile	Regulatory	Enterprise + extended L6 obligations	None beyond undeclared extension fields
Extension	SOC_SIEM	atlas_technique_id and approved security-tooling mapping fields	Requires extensions to contain SOC_SIEM
Extension	AI_Security	Selected L5 fields: agent_escalation_chain, mcp_endpoint_ids, oauth_pivot_chain, policy_enforcement_point, ai_system_type, non_adversarial_failure_type	Requires extensions to contain AI_Security

Core Profile implementations SHALL NOT include prohibited L3/L4 fields in the primary payload. Extension-gated fields SHALL NOT appear unless the corresponding extension is declared. Extended data MAY be stored under x_extensions where permitted, but SHALL NOT affect base-profile validation. Custom vocabulary values SHALL use the x_ prefix only where the controlled vocabulary explicitly allows it.

#	Conformance Requirement	Validation Method
1	Mandatory fields for selected profile are present	Schema validation
2	Mandatory fields contain valid type, format, enum, and normative precision values	Schema + semantic precision validation
3	Prohibited fields are absent from selected base profile primary payload	Field-presence validation
4	Extension-gated fields occur only when the matching extension is declared	Schema allOf validation
5	Severity calculations match the normative reference implementation	Deterministic output comparison
6	Conditional validation rules are enforced	Rule-based validation
7	severity_rationale includes calibration_status and required provisional notices	Semantic check
8	regulatory_trigger presence requires trigger_confidence and human_review_recommended	Schema allOf validation
9	root_cause_specific_type uses controlled vocabulary or x_ extension value	Vocabulary check

4. Severity Scoring Engine v1.0

UAIF severity scoring draws structural inspiration from CVSS v3.1/v4.0 but is a post-incident classification and routing model. It is not a replacement for vulnerability scoring.

Harm Type	Default Weight	Status
Physical_Harm	4.0	Provisional pending sector validation
Environmental_Harm	3.5	Provisional pending sector validation
Security_Integrity	3.5	Provisional pending sector validation
Privacy_Violation	3.0	Provisional pending sector validation
Financial_Loss	2.5	Provisional pending sector validation
Psychological_Harm	2.0	Provisional pending sector validation
Property_Damage	2.0	Provisional pending sector validation
Reputational_Harm	1.5	Provisional pending sector validation

4.1 Context Modifiers

Context	Base Value	Effective Cap
Model_Poisoning	2.0	1.5 per factor
RAG_leakage	1.5	1.5
Prompt_Injection	1.5	1.5
Agent_Hijack	1.5	1.5
Agentic	1.0	1.0
Adversarial	0.5	0.5

4.2 Normative Calculation

The normative UAIF v1.0 calculation SHALL follow this order: dominant-harm anchor; diminishing context contribution; impact scale using MAX(population, blast radius); inverted temporal confidence reduction; final analytical and presentation scores. Implementations SHALL NOT use collapsed single-score fields or non-normative stakeholder multipliers.

```
from decimal import Decimal, ROUND_HALF_UP

HARM_WEIGHTS = {
    "Physical_Harm": 4.0, "Environmental_Harm": 3.5, "Security_Integrity": 3.5,
    "Privacy_Violation": 3.0, "Financial_Loss": 2.5, "Psychological_Harm": 2.0,
    "Property_Damage": 2.0, "Reputational_Harm": 1.5
}

CONTEXT_MODIFIERS = {
    "Model_Poisoning": 2.0, "RAG_leakage": 1.5, "Prompt_Injection": 1.5,
    "Agent_Hijack": 1.5, "Agentic": 1.0, "Adversarial": 0.5
}

def _round_decimal(value, places=1):
    quantum = Decimal("1." + ("0" * places)) if places else Decimal("1")
    return float(Decimal(str(value)).quantize(quantum, rounding=ROUND_HALF_UP))

def score_to_level(score):
    if score < 2.0: return 1
    if score < 4.0: return 2
    if score < 6.5: return 3
    if score < 8.5: return 4
    return 5

# Reference implementation applies dominant-harm weighting, diminishing context contribution,
# impact scale, confidence/temporal adjustment, and returns severity_rationale as the canonical audit object.
```

4.3 Impact Scale and Temporal Adjustment

Population	Multiplier	Blast Radius	Multiplier
1-10	0.8	single_system	1.0
11-100	1.0	business_unit	1.2
101-1,000	1.2	entire_org	1.4
1,001-10,000	1.4	cross_org	1.6
>10,000	1.6	critical_infra	1.8

Impact scale = min(1.8, max(population_multiplier, blast_radius_multiplier)). This prevents double-counting of scale dimensions.

4.4 Rounding and Precision Rule

All decimal scores SHALL use Decimal ROUND_HALF_UP. severity_analytical_score and severity_presentation_score SHALL be represented to one decimal place as a normative semantic/conformance rule. Ratio or internal contribution values MAY be rounded to two decimal places where specified by severity_rationale.

The canonical JSON Schema intentionally does not use multipleOf: 0.1 for score precision. Binary floating-point implementations can produce validator-dependent false negatives for valid decimal values such as 2.8, 5.1, or 9.1. Precision SHALL therefore be verified by the reference implementation or equivalent conformance tooling rather than by JSON Schema multipleOf.

4.5 Severity Matrix

Level	Presentation Score Range	Indicative Harm Threshold	Regulatory Routing
1 Info	0.0 <= score < 2.0	No external harm	None

Level	Presentation Score Range	Indicative Harm Threshold	Regulatory Routing
2 Low	2.0 <= score < 4.0	Minor confusion or < \$10K impact	Internal audit only
3 Medium	4.0 <= score < 6.5	Reputational/non-sensitive data/\$10K-\$100K	Sectoral reporting if applicable
4 High	6.5 <= score < 8.5	Physical safety, sensitive data, >\$100K	EU AI Act Article 73 technical proxy where applicable
5 Critical	8.5 <= score <= 10.0	Widespread societal harm, critical infrastructure, major systemic impact	EU AI Act Article 73 serious-incident proxy where applicable

Boundary rule: lower bounds are inclusive; upper bounds are exclusive except Level 5, where 10.0 is inclusive.

5. Regulatory Trigger Logic

UAIF regulatory_trigger is a technical routing flag. It SHALL NOT be treated as a legal determination. EU AI Act serious-incident references SHALL use Article 73. Obsolete EU AI Act incident-reporting proxy references SHALL NOT be used in UAIF v1.0.

Parameter	Normative Rule
threshold_score	Default 6.5, evaluated against severity_presentation_score.
regulated_zones	Default EEA jurisdictions: AT, BE, BG, HR, CY, CZ, DK, EE, FI, FR, DE, GR, HU, IE, IT, LV, LT, LU, MT, NL, PL, PT, RO, SK, SI, ES, SE, plus NO, IS, LI. Implementations MAY override by policy.
confidence_bands	sector_validated: HIGH if distance >=0.3 else MEDIUM. provisional: MEDIUM if distance >=0.5 else LOW. custom: MEDIUM.
human_review	Required when trigger_confidence is LOW; recommended where legal obligations are uncertain.
reporting_timelines	EU AI Act Article 73 timing must be assessed by legal counsel. GDPR/NIS2/DORA timelines are separate and SHALL NOT be copied into EU AI Act routing logic.

6. Sector Taxonomy

The normative sector_code field SHALL use the 19-code NIS2-aligned UAIF enum below. NAICS/NACE identifiers MAY be captured in optional mapping metadata fields such as nace_code_actual or naics_code_actual; they SHALL NOT replace sector_code.

sector_code	Usage
ENERGY	Normative UAIF sector value
TRANSPORT	Normative UAIF sector value
BANKING	Normative UAIF sector value
FINANCIAL_MARKETS	Normative UAIF sector value
HEALTH	Normative UAIF sector value
DRINKING_WATER	Normative UAIF sector value
WASTE_WATER	Normative UAIF sector value
DIGITAL_INFRA	Normative UAIF sector value
ICT_SERVICES	Normative UAIF sector value
PUBLIC_ADMIN	Normative UAIF sector value
SPACE	Normative UAIF sector value
POSTAL	Normative UAIF sector value
WASTE_MGMT	Normative UAIF sector value
CHEMICALS	Normative UAIF sector value
FOOD	Normative UAIF sector value
MANUFACTURING	Normative UAIF sector value

sector_code	Usage
TELECOM	Normative UAIF sector value
RESEARCH	Normative UAIF sector value
OTHER	Normative UAIF sector value

7. Governance, Versioning, and Supporting Documents

Item	Rule
Version display	Use UAIF™ v1.0 in titles and market-facing text. Machine-readable schema metadata MAY use 1.0.0 only where explicitly described as semantic-version metadata.
Final Publication status	Normative architecture and conformance behavior are frozen for v1.0. Errata may correct non-behavioral defects; substantive normative behavior changes require a future revision.
Change control	ODA3 Institute maintains sole versioning and interpretation authority for UAIF v1.0.
Supporting Document D2	Contributor Licence Agreement; separate companion document, not included in this specification.
Supporting Document D4	Data Contribution Agreement; separate companion document, not included in this specification.
HMAC key management	Deduplication hash keys SHOULD be generated and stored under organisational KMS/HSM policy, rotated at least annually or after suspected compromise, and segregated by environment.
GEL references	Commercial use, embedding, certification claims, and conformance programme operation are governed by GEL v1.0.

8. Licensing and Brand Use

UAIF™, GAISSE™, AI-IRF™, ODA3™, and ODA3 Institute marks are proprietary marks of ODA3 Pvt Ltd. This document is licensed under GEL v1.0. No certification, endorsement, or official conformance status is granted by use of this document or schema alone.

Appendix A: UAIF Core Profile Schema Reference

The normative machine-readable schema is maintained in TCR-STD-003. The schema identifier is:
<https://schemas.oda3.org/uaif/core/v1.0/uaif-core-v1.0.json>

Appendix B: Conditional Validation Rules

#	Condition	Consequence
1	root_cause_specific_type = Adversarial_Attack	exploit_path_description SHALL be present
2	record_type = Incident (any profile)	severity_analytical_score, severity_presentation_score, severity_level, severity_rationale SHALL be present
2a	record_type = Incident AND profile ∈ {Enterprise, Regulatory}	realized_harm_categories SHALL additionally be present
3	regulatory_trigger present	trigger_confidence and human_review_recommended SHALL be present
4	severity_rationale.calibration_status = provisional	weight_provisional_notice SHALL be present in severity_rationale
5	context_modifiers_applied contains entries	Each entry SHALL include context, base_value, contribution_factor, effective_contribution
6	root_cause_specific_type not in controlled vocabulary	Value SHALL be prefixed x_ or validation fails

#	Condition	Consequence
7	profile = Core	L3/L4 fields SHALL be absent from primary payload
8	atlas_technique_id present	extensions SHALL contain SOC_SIEM
9	AI Security L5 field present	extensions SHALL contain AI_Security

Appendix H: Harm Weight Derivation Methodology

Default harm weights are provisional and derived from preliminary incident review and expert elicitation. Target inter-rater reliability for sector validation is ICC ≥ 0.80 using ICC(2,k), two-way mixed-effects model, absolute agreement, average measures. Sector validation requires ≥ 10 raters and ≥ 30 reference incidents per sector cohort.

Appendix I: Standards Alignment Notes

NIST AI RMF 1.0 is the current alignment basis for GOVERN, MAP, MEASURE, and MANAGE functions. NIST AI 100-2 E2025 may be referenced only for adversarial machine learning taxonomy alignment and not as a successor AI RMF release.

EU AI Act references in UAIF v1.0 use Article 73 for serious incident reporting proxies. Legal counsel determines applicability and reporting timelines.

Appendix J: NACE and NAICS Mapping Guidance

NACE and NAICS codes are mapping metadata only. Implementations SHOULD map primary activity codes to the normative UAIF sector_code enum and retain the original activity code in nace_code_actual or naics_code_actual.

Appendix K: Controlled Vocabulary

Vocabulary	Values
record_type	Incident, Vulnerability
profile	Core, Enterprise, Regulatory
extensions	SOC_SIEM, AI_Security
root_cause_specific_type	Adversarial_Attack, Prompt_Injection, Model_Poisoning, Data_Poisoning, RAG_Leakage, Hallucination, Agent_Escalation, Policy_Violation, System_Failure, Human_Error, Supply_Chain_Compromise, Unauthorised_Access, Data_Breach, Denial_of_Service, Misuse, x_[custom]
realized_harm_categories	Physical_Harm, Environmental_Harm, Security_Integrity, Privacy_Violation, Financial_Loss, Psychological_Harm, Property_Damage, Reputational_Harm