

ODA3 Institute

GAISSF Energy Sector Implementation Guide

Operational interpretation of GAISSF v1.0 for electricity, oil and gas, storage, renewables and energy-market environments

Field	Value
Document ID	SEC-044
Version	1.0
Classification	Informative sector implementation guidance
Status	Publication Candidate - specialist approvals open
Publisher	ODA3 Institute
Publication date	30 June 2026
Authoritative source	GAISSF-NOR-001 v1.0
Control baseline	59 controls across D1-D9
Publication channel	Website / GitHub

© 2026 ODA3 Pvt Ltd. All rights reserved. Published by ODA3 Institute.

This document does not constitute legal, regulatory, engineering, safety, audit or certification advice and does not guarantee security, safety, reliability, compliance or absence of harmful outcomes.

Document control and precedence

GAISSF-NOR-001 is the authoritative normative source. This guide is informative and does not amend, replace, narrow or create GAISSF requirements. The F/O/A sequencing labels in this guide are implementation aids, not GAISSF conformance tiers and not universal priorities.

Executive overview

Energy-sector AI can influence planning, maintenance, cybersecurity, markets, restoration and physical operation. Risk increases materially when AI output can affect commands, operating envelopes, safety functions, grid or process stability, essential customer service or emergency restoration. The implementation objective is to bound authority, preserve independent protection, test realistic failure modes and retain reconstructable evidence.

- Accountability remains with the energy organisation even where models, cloud services, data or support are outsourced.
- Recommendation, approval, scheduling, command, physical actuation and verified outcome are separate control states.
- Model accuracy is not evidence of operational safety; degraded-mode, communications-loss, fallback and recovery require validation.
- Human presence is not effective oversight without competence, information, time, authority, override and escalation.
- Crosswalks support analysis but do not prove legal or regulatory equivalence.

Scope, exclusions and applicability

The guide covers electricity generation, transmission, distribution, markets and retail; oil and gas; renewable generation; energy storage; hydrogen; field operations; customer systems; enterprise functions and supporting cybersecurity. Nuclear applications require specialist nuclear safety, security and regulatory review. Aviation assurance references are contextual only where airborne systems are actually in scope.

Not Applicable criteria

- A control may be marked Not Applicable only where the system boundary is explicit and evidence shows the addressed risk is absent.
- The rationale must be approved by an accountable owner and include a reassessment trigger.
- Technical difficulty, cost, supplier refusal or low maturity does not make a control Not Applicable; those conditions require treatment, compensation or accepted residual risk.
- Example: a purely analytical load-forecast model with no actuator, SCADA or command interface may justify non-applicability of a direct physical-actuation control, subject to documented verification.

Operational Technology and AI: a primer

Concept	Energy-sector meaning
IT and OT objectives	IT commonly emphasises confidentiality, integrity and availability. OT often orders priorities as availability, integrity and confidentiality, while safety and resilience remain separate critical outcomes.
Deterministic control vs probabilistic AI	Control and protection functions require predictable bounded behaviour. AI output is probabilistic and should not silently replace deterministic constraints, interlocks or protection logic.
Purdue-style layering	AI may appear in enterprise, operations-management, supervisory, control or edge layers. Trust boundaries and

Concept	Energy-sector meaning
	permitted data/command flows must be explicit; the model is a reference pattern, not a mandatory architecture.
Independent Protection Layer	A protection layer reduces risk independently of the initiating event and the AI path. AI must not be the sole protection against a hazard it can create or fail to detect.
Safety instrumented/protection systems	Safety and protection functions require specialist engineering and lifecycle assurance. AI should not bypass or weaken their independence, integrity or safe-state behaviour.

Energy AI system taxonomy

Class	Role	Authority	Assurance expectation
1	Enterprise support	No operational authority	Standard governance, security and evidence.
2	Operational advisory	Recommendation only	Representative validation and operator review.
3	Operational decision support	Influences material decision	Independent validation, uncertainty display and tested fallback.
4	High-impact automated	Executes bounded action/workflow	Machine identity, action constraints, monitoring, rollback and stronger assurance.
5	Closed-loop or safety-relevant	Direct physical or safety consequence	Independent protection, safe state, override, simulation/HIL where feasible, and ordinarily A4-A5 assurance.

Energy risk model: Confidentiality, Integrity, Availability, Safety and Resilience

Dimension	Definition in energy AI	Representative questions
Confidentiality	Protection of sensitive topology, asset, vulnerability, customer, commercial and operational data.	Could disclosure enable physical, cyber, privacy, market or competitive harm?
Integrity	Confidence that data, model, configuration, recommendation, command and record are complete, authentic and unmanipulated.	Can the organisation detect poisoned data, altered topology, substituted models or modified commands?
Availability	Ability to provide required capability within operational time constraints without unsafe dependence.	What happens during cloud, telecom, data-source, model-service or identity-system loss?
Safety	Prevention of unacceptable harm to people, equipment, environment and critical operations.	Can an AI-induced action exceed an operating envelope, defeat protection or delay emergency action?
Resilience	Ability to degrade gracefully, retain safe or essential operation, recover trusted state and learn from events.	Are fallback, rollback, restoration, evidence preservation and post-event reassessment tested?

Assess consequence, likelihood, exposure, detectability, reversibility and confidence separately. Do not collapse safety or resilience into a single generic cybersecurity score.

Implementation sequencing aid

F - Foundational identifies controls usually needed to establish scope, authority, evidence and basic containment. O - Operational identifies controls generally applied as systems are deployed and operated. A - Advanced/contextual

identifies controls whose applicability or implementation depth is more technology- or risk-specific. These labels do not change normative applicability.

Label	Purpose	Caution
F - Foundational	Start with inventory, accountability, auditability, authority and core integrity/monitoring.	Foundational does not mean sufficient for high-consequence operation.
O - Operational	Implement, test and operate controls proportionate to system risk and authority.	Operational controls may be required from day one for a particular system.
A - Advanced/contextual	Address specialised technologies, threats or higher-assurance contexts.	Advanced does not mean optional when the underlying risk is present.

Control-by-control energy interpretations

D1-CTL-01 - DATASET PROVENANCE & POISONING PREVENTION

Field	Energy-sector treatment
Domain	D1: MODEL INTEGRITY & ADVERSARIAL ROBUSTNESS
Implementation sequence	F - Foundational
Normative reference	GAISSF-NOR-001, D1-CTL-01
Interpretation	In an energy environment, dataset provenance & poisoning prevention is implemented by applying verify provenance, representativeness, topology and telemetry integrity, drift, poisoning resistance, version integrity and reproducibility. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use signed or otherwise verifiable lineage; validate seasonal and post-reconfiguration conditions; test stale, missing and manipulated telemetry; validate topology and time synchronisation.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Lineage/provenance record; representative-data assessment; integrity or drift test report; version and approval record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is dataset provenance & poisoning prevention scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D1-CTL-02 - MODEL EXTRACTION RESISTANCE

Field	Energy-sector treatment
Domain	D1: MODEL INTEGRITY & ADVERSARIAL ROBUSTNESS
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D1-CTL-02
Interpretation	In an energy environment, model extraction resistance is implemented by applying verify provenance, representativeness, topology and telemetry integrity, drift, poisoning resistance, version integrity and reproducibility. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use signed or otherwise verifiable lineage; validate seasonal and post-reconfiguration conditions; test stale, missing and manipulated telemetry; validate topology and time synchronisation.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Lineage/provenance record; representative-data assessment; integrity or drift test report; version and approval record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is model extraction resistance scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D1-CTL-03 - BEHAVIORAL DRIFT DETECTION

Field	Energy-sector treatment
Domain	D1: MODEL INTEGRITY & ADVERSARIAL ROBUSTNESS
Implementation sequence	F - Foundational
Normative reference	GAISSF-NOR-001, D1-CTL-03
Interpretation	In an energy environment, behavioral drift detection is implemented by applying verify provenance, representativeness, topology and telemetry integrity, drift, poisoning resistance, version integrity and reproducibility. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use signed or otherwise verifiable lineage; validate seasonal and post-reconfiguration conditions; test stale, missing and manipulated telemetry; validate topology and time synchronisation.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned

Field	Energy-sector treatment
	operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Lineage/provenance record; representative-data assessment; integrity or drift test report; version and approval record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is behavioral drift detection scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D1-CTL-04 - FEDERATED LEARNING POISONING PREVENTION

Field	Energy-sector treatment
Domain	D1: MODEL INTEGRITY & ADVERSARIAL ROBUSTNESS
Implementation sequence	O - Operational
Normative reference	GAISF-NOR-001, D1-CTL-04
Interpretation	In an energy environment, federated learning poisoning prevention is implemented by applying verify provenance, representativeness, topology and telemetry integrity, drift, poisoning resistance, version integrity and reproducibility. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use signed or otherwise verifiable lineage; validate seasonal and post-reconfiguration conditions; test stale, missing and manipulated telemetry; validate topology and time synchronisation.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Lineage/provenance record; representative-data assessment; integrity or drift test report; version and approval record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is federated learning poisoning prevention scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.

Field	Energy-sector treatment
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D1-CTL-05 - EMBEDDING SPACE ROBUSTNESS

Field	Energy-sector treatment
Domain	D1: MODEL INTEGRITY & ADVERSARIAL ROBUSTNESS
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D1-CTL-05
Interpretation	In an energy environment, embedding space robustness is implemented by applying verify provenance, representativeness, topology and telemetry integrity, drift, poisoning resistance, version integrity and reproducibility. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use signed or otherwise verifiable lineage; validate seasonal and post-reconfiguration conditions; test stale, missing and manipulated telemetry; validate topology and time synchronisation.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Lineage/provenance record; representative-data assessment; integrity or drift test report; version and approval record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is embedding space robustness scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D1-CTL-06 - POST-QUANTUM MODEL SIGNING & CRYPTO HARDENING

Field	Energy-sector treatment
Domain	D1: MODEL INTEGRITY & ADVERSARIAL ROBUSTNESS
Implementation sequence	A - Advanced / contextual
Normative reference	GAISSF-NOR-001, D1-CTL-06
Interpretation	In an energy environment, post-quantum model signing & crypto hardening is implemented by applying verify provenance, representativeness, topology and telemetry integrity, drift, poisoning resistance, version integrity and reproducibility. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative

Field	Energy-sector treatment
	testing; operational monitoring; reconstructable evidence.
Recommended practices	Use signed or otherwise verifiable lineage; validate seasonal and post-reconfiguration conditions; test stale, missing and manipulated telemetry; validate topology and time synchronisation.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Lineage/provenance record; representative-data assessment; integrity or drift test report; version and approval record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is post-quantum model signing & crypto hardening scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D1-CTL-07 - LORA/ADAPTER INTEGRITY VERIFICATION

Field	Energy-sector treatment
Domain	D1: MODEL INTEGRITY & ADVERSARIAL ROBUSTNESS
Implementation sequence	O - Operational
Normative reference	GAISF-NOR-001, D1-CTL-07
Interpretation	In an energy environment, lora/adaptor integrity verification is implemented by applying verify provenance, representativeness, topology and telemetry integrity, drift, poisoning resistance, version integrity and reproducibility. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use signed or otherwise verifiable lineage; validate seasonal and post-reconfiguration conditions; test stale, missing and manipulated telemetry; validate topology and time synchronisation.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Lineage/provenance record; representative-data assessment; integrity or drift test report; version and approval record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is lora/adaptor integrity verification scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?

Field	Energy-sector treatment
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D1-CTL-08 - MODEL MERGE ATTACK DETECTION

Field	Energy-sector treatment
Domain	D1: MODEL INTEGRITY & ADVERSARIAL ROBUSTNESS
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D1-CTL-08
Interpretation	In an energy environment, model merge attack detection is implemented by applying verify provenance, representativeness, topology and telemetry integrity, drift, poisoning resistance, version integrity and reproducibility. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use signed or otherwise verifiable lineage; validate seasonal and post-reconfiguration conditions; test stale, missing and manipulated telemetry; validate topology and time synchronisation.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Lineage/provenance record; representative-data assessment; integrity or drift test report; version and approval record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is model merge attack detection scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D1-CTL-09 - QUANTIZATION BACKDOOR SCREENING

Field	Energy-sector treatment
Domain	D1: MODEL INTEGRITY & ADVERSARIAL ROBUSTNESS
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D1-CTL-09
Interpretation	In an energy environment, quantization backdoor screening is implemented by applying verify provenance, representativeness, topology and telemetry integrity, drift, poisoning resistance, version integrity and reproducibility. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an

Field	Energy-sector treatment
	accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use signed or otherwise verifiable lineage; validate seasonal and post-reconfiguration conditions; test stale, missing and manipulated telemetry; validate topology and time synchronisation.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Lineage/provenance record; representative-data assessment; integrity or drift test report; version and approval record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is quantization backdoor screening scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D2-CTL-01 - DIRECT PROMPT INJECTION PREVENTION

Field	Energy-sector treatment
Domain	D2: RUNTIME SECURITY & ADVERSARIAL DEFENSE
Implementation sequence	F - Foundational
Normative reference	GAISSE-NOR-001, D2-CTL-01
Interpretation	In an energy environment, direct prompt injection prevention is implemented by applying separate untrusted content from trusted operational instructions; constrain interfaces, retrieval, multimodal inputs and tool calls. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use schema validation, trust separation, interface allow-lists, rate limits and monitoring for rejected or anomalous requests; prevent external content from becoming operational instruction.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Threat model; interface and input-control configuration; adversarial test report; runtime alert and response record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings;

Field	Energy-sector treatment
	fallback invocations; incidents and near misses.
Assessment question	How is direct prompt injection prevention scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D2-CTL-02 - INDIRECT PROMPT INJECTION PREVENTION

Field	Energy-sector treatment
Domain	D2: RUNTIME SECURITY & ADVERSARIAL DEFENSE
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D2-CTL-02
Interpretation	In an energy environment, indirect prompt injection prevention is implemented by applying separate untrusted content from trusted operational instructions; constrain interfaces, retrieval, multimodal inputs and tool calls. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use schema validation, trust separation, interface allow-lists, rate limits and monitoring for rejected or anomalous requests; prevent external content from becoming operational instruction.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Threat model; interface and input-control configuration; adversarial test report; runtime alert and response record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is indirect prompt injection prevention scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D2-CTL-03 - JAILBREAK RESISTANCE TESTING

Field	Energy-sector treatment
Domain	D2: RUNTIME SECURITY & ADVERSARIAL DEFENSE
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D2-CTL-03
Interpretation	In an energy environment, jailbreak resistance testing is implemented by applying separate untrusted content from trusted operational instructions; constrain interfaces, retrieval, multimodal inputs and tool calls. Implementation intensity is determined by operational authority,

Field	Energy-sector treatment
	safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use schema validation, trust separation, interface allow-lists, rate limits and monitoring for rejected or anomalous requests; prevent external content from becoming operational instruction.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Threat model; interface and input-control configuration; adversarial test report; runtime alert and response record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is jailbreak resistance testing scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D2-CTL-04 - MULTI-MODAL INJECTION DEFENSE

Field	Energy-sector treatment
Domain	D2: RUNTIME SECURITY & ADVERSARIAL DEFENSE
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D2-CTL-04
Interpretation	In an energy environment, multi-modal injection defense is implemented by applying separate untrusted content from trusted operational instructions; constrain interfaces, retrieval, multimodal inputs and tool calls. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use schema validation, trust separation, interface allow-lists, rate limits and monitoring for rejected or anomalous requests; prevent external content from becoming operational instruction.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Threat model; interface and input-control configuration; adversarial test

Field	Energy-sector treatment
	report; runtime alert and response record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is multi-modal injection defense scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D2-CTL-05 - FUNCTION CALL/TOOL CALL INJECTION PREVENTION

Field	Energy-sector treatment
Domain	D2: RUNTIME SECURITY & ADVERSARIAL DEFENSE
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D2-CTL-05
Interpretation	In an energy environment, function call/tool call injection prevention is implemented by applying separate untrusted content from trusted operational instructions; constrain interfaces, retrieval, multimodal inputs and tool calls. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use schema validation, trust separation, interface allow-lists, rate limits and monitoring for rejected or anomalous requests; prevent external content from becoming operational instruction.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Threat model; interface and input-control configuration; adversarial test report; runtime alert and response record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is function call/tool call injection prevention scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D2-CTL-06 - CROSS-CONTEXT HIJACKING MITIGATION

Field	Energy-sector treatment
Domain	D2: RUNTIME SECURITY & ADVERSARIAL DEFENSE
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D2-CTL-06

Field	Energy-sector treatment
Interpretation	In an energy environment, cross-context hijacking mitigation is implemented by applying separate untrusted content from trusted operational instructions; constrain interfaces, retrieval, multimodal inputs and tool calls. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use schema validation, trust separation, interface allow-lists, rate limits and monitoring for rejected or anomalous requests; prevent external content from becoming operational instruction.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Threat model; interface and input-control configuration; adversarial test report; runtime alert and response record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is cross-context hijacking mitigation scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D3-CTL-01 - LEAST AGENCY ENFORCEMENT

Field	Energy-sector treatment
Domain	D3: AGENTIC RISK & AUTONOMOUS SYSTEM SECURITY
Implementation sequence	F - Foundational
Normative reference	GAISSF-NOR-001, D3-CTL-01
Interpretation	In an energy environment, least agency enforcement is implemented by applying separate recommendation, approval, scheduling, command and actuation; apply least agency, machine identity and bounded actions. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use machine identity, least agency, dual authorisation for material actions, action envelopes, rate-of-change constraints and post-action reconciliation.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.

Field	Energy-sector treatment
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Authority matrix; machine-identity record; action-policy configuration; approval, override and command-reconciliation logs. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is least agency enforcement scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D3-CTL-02 - INTER-AGENT COMMUNICATION SECURITY

Field	Energy-sector treatment
Domain	D3: AGENTIC RISK & AUTONOMOUS SYSTEM SECURITY
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D3-CTL-02
Interpretation	In an energy environment, inter-agent communication security is implemented by applying separate recommendation, approval, scheduling, command and actuation; apply least agency, machine identity and bounded actions. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use machine identity, least agency, dual authorisation for material actions, action envelopes, rate-of-change constraints and post-action reconciliation.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Authority matrix; machine-identity record; action-policy configuration; approval, override and command-reconciliation logs. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is inter-agent communication security scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D3-CTL-03 - AGENTIC PROMPT CHAINING DETECTION

Field	Energy-sector treatment
Domain	D3: AGENTIC RISK & AUTONOMOUS SYSTEM SECURITY
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D3-CTL-03
Interpretation	In an energy environment, agentic prompt chaining detection is implemented by applying separate recommendation, approval, scheduling, command and actuation; apply least agency, machine identity and bounded actions. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use machine identity, least agency, dual authorisation for material actions, action envelopes, rate-of-change constraints and post-action reconciliation.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Authority matrix; machine-identity record; action-policy configuration; approval, override and command-reconciliation logs. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is agentic prompt chaining detection scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D3-CTL-04 - EMBODIED AI SAFETY CONTROLS

Field	Energy-sector treatment
Domain	D3: AGENTIC RISK & AUTONOMOUS SYSTEM SECURITY
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D3-CTL-04
Interpretation	In an energy environment, embodied ai safety controls is implemented by applying separate recommendation, approval, scheduling, command and actuation; apply least agency, machine identity and bounded actions. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use machine identity, least agency, dual authorisation for material actions, action envelopes, rate-of-change constraints and post-action reconciliation.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned

Field	Energy-sector treatment
	operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Authority matrix; machine-identity record; action-policy configuration; approval, override and command-reconciliation logs. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is embodied ai safety controls scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D3-CTL-05 - MULTI-AGENT TRUST CHAIN ATTESTATION

Field	Energy-sector treatment
Domain	D3: AGENTIC RISK & AUTONOMOUS SYSTEM SECURITY
Implementation sequence	O - Operational
Normative reference	GAISF-NOR-001, D3-CTL-05
Interpretation	In an energy environment, multi-agent trust chain attestation is implemented by applying separate recommendation, approval, scheduling, command and actuation; apply least agency, machine identity and bounded actions. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use machine identity, least agency, dual authorisation for material actions, action envelopes, rate-of-change constraints and post-action reconciliation.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Authority matrix; machine-identity record; action-policy configuration; approval, override and command-reconciliation logs. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is multi-agent trust chain attestation scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.

Field	Energy-sector treatment
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D3-CTL-06 - PERSISTENT MEMORY EXFILTRATION PREVENTION

Field	Energy-sector treatment
Domain	D3: AGENTIC RISK & AUTONOMOUS SYSTEM SECURITY
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D3-CTL-06
Interpretation	In an energy environment, persistent memory exfiltration prevention is implemented by applying separate recommendation, approval, scheduling, command and actuation; apply least agency, machine identity and bounded actions. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use machine identity, least agency, dual authorisation for material actions, action envelopes, rate-of-change constraints and post-action reconciliation.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Authority matrix; machine-identity record; action-policy configuration; approval, override and command-reconciliation logs. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is persistent memory exfiltration prevention scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D3-CTL-07 - SECURE MEMORY LIFECYCLE MANAGEMENT

Field	Energy-sector treatment
Domain	D3: AGENTIC RISK & AUTONOMOUS SYSTEM SECURITY
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D3-CTL-07
Interpretation	In an energy environment, secure memory lifecycle management is implemented by applying separate recommendation, approval, scheduling, command and actuation; apply least agency, machine identity and bounded actions. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.

Field	Energy-sector treatment
Recommended practices	Use machine identity, least agency, dual authorisation for material actions, action envelopes, rate-of-change constraints and post-action reconciliation.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Authority matrix; machine-identity record; action-policy configuration; approval, override and command-reconciliation logs. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is secure memory lifecycle management scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D4-CTL-01 - AI BILL OF MATERIALS (AI BOM) MAINTENANCE

Field	Energy-sector treatment
Domain	D4: SUPPLY CHAIN & THIRD-PARTY AI SECURITY
Implementation sequence	F - Foundational
Normative reference	GAISSF-NOR-001, D4-CTL-01
Interpretation	In an energy environment, ai bill of materials (ai bom) maintenance is implemented by applying assess model, component, data, cloud, integrator and remote-support provenance, update security, continuity and concentration risk. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Require provenance, secure update and rollback, vulnerability and incident terms, subcontractor transparency, continuity tests and exit arrangements.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Supplier due-diligence record; AI/component inventory; provenance attestation; update validation; continuity and exit test. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is ai bill of materials (ai bom) maintenance scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability;

Field	Energy-sector treatment
	missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D4-CTL-02 - MODEL FILE & ARTIFACT SCANNING

Field	Energy-sector treatment
Domain	D4: SUPPLY CHAIN & THIRD-PARTY AI SECURITY
Implementation sequence	F - Foundational
Normative reference	GAISSF-NOR-001, D4-CTL-02
Interpretation	In an energy environment, model file & artifact scanning is implemented by applying assess model, component, data, cloud, integrator and remote-support provenance, update security, continuity and concentration risk. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Require provenance, secure update and rollback, vulnerability and incident terms, subcontractor transparency, continuity tests and exit arrangements.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Supplier due-diligence record; AI/component inventory; provenance attestation; update validation; continuity and exit test. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is model file & artifact scanning scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D4-CTL-03 - MODEL HUB & REGISTRY VETTING

Field	Energy-sector treatment
Domain	D4: SUPPLY CHAIN & THIRD-PARTY AI SECURITY
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D4-CTL-03
Interpretation	In an energy environment, model hub & registry vetting is implemented by applying assess model, component, data, cloud, integrator and remote-support provenance, update security, continuity and concentration risk. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty

Field	Energy-sector treatment
	alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Require provenance, secure update and rollback, vulnerability and incident terms, subcontractor transparency, continuity tests and exit arrangements.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Supplier due-diligence record; AI/component inventory; provenance attestation; update validation; continuity and exit test. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is model hub & registry vetting scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D4-CTL-04 - MCP SERVER BEHAVIORAL MONITORING

Field	Energy-sector treatment
Domain	D4: SUPPLY CHAIN & THIRD-PARTY AI SECURITY
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D4-CTL-04
Interpretation	In an energy environment, mcp server behavioral monitoring is implemented by applying assess model, component, data, cloud, integrator and remote-support provenance, update security, continuity and concentration risk. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Require provenance, secure update and rollback, vulnerability and incident terms, subcontractor transparency, continuity tests and exit arrangements.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Supplier due-diligence record; AI/component inventory; provenance attestation; update validation; continuity and exit test. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.

Field	Energy-sector treatment
Assessment question	How is mcp server behavioral monitoring scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D4-CTL-05 - THIRD-PARTY AI API SECURITY ASSESSMENT

Field	Energy-sector treatment
Domain	D4: SUPPLY CHAIN & THIRD-PARTY AI SECURITY
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D4-CTL-05
Interpretation	In an energy environment, third-party ai api security assessment is implemented by applying assess model, component, data, cloud, integrator and remote-support provenance, update security, continuity and concentration risk. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Require provenance, secure update and rollback, vulnerability and incident terms, subcontractor transparency, continuity tests and exit arrangements.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Supplier due-diligence record; AI/component inventory; provenance attestation; update validation; continuity and exit test. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is third-party ai api security assessment scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D4-CTL-06 - SHADOW AI DISCOVERY & GOVERNANCE

Field	Energy-sector treatment
Domain	D4: SUPPLY CHAIN & THIRD-PARTY AI SECURITY
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D4-CTL-06
Interpretation	In an energy environment, shadow ai discovery & governance is implemented by applying assess model, component, data, cloud, integrator and remote-support provenance, update security, continuity and concentration risk. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality,

Field	Energy-sector treatment
	connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Require provenance, secure update and rollback, vulnerability and incident terms, subcontractor transparency, continuity tests and exit arrangements.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Supplier due-diligence record; AI/component inventory; provenance attestation; update validation; continuity and exit test. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is shadow ai discovery & governance scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D4-CTL-07 - AI SOFTWARE COMPOSITION ANALYSIS (SCA)

Field	Energy-sector treatment
Domain	D4: SUPPLY CHAIN & THIRD-PARTY AI SECURITY
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D4-CTL-07
Interpretation	In an energy environment, ai software composition analysis (sca) is implemented by applying assess model, component, data, cloud, integrator and remote-support provenance, update security, continuity and concentration risk. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Require provenance, secure update and rollback, vulnerability and incident terms, subcontractor transparency, continuity tests and exit arrangements.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Supplier due-diligence record; AI/component inventory; provenance attestation; update validation; continuity and exit test. Retention follows

Field	Energy-sector treatment
	applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is ai software composition analysis (sca) scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D5-CTL-01 - HARMFUL CONTENT BLOCKING

Field	Energy-sector treatment
Domain	D5: CONTENT SAFETY & OUTPUT INTEGRITY
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D5-CTL-01
Interpretation	In an energy environment, harmful content blocking is implemented by applying test outputs for unsafe, deceptive, confidential or operationally misleading content; communicate uncertainty and freshness. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Display confidence, uncertainty, freshness and operational limitations; prevent unreviewed output from becoming a command, authoritative record or safety conclusion.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Output evaluation; safety/content test record; uncertainty-display evidence; exception and escalation record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is harmful content blocking scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D5-CTL-02 - PII LEAKAGE PREVENTION

Field	Energy-sector treatment
Domain	D5: CONTENT SAFETY & OUTPUT INTEGRITY
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D5-CTL-02
Interpretation	In an energy environment, pii leakage prevention is implemented by

Field	Energy-sector treatment
	applying test outputs for unsafe, deceptive, confidential or operationally misleading content; communicate uncertainty and freshness. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Display confidence, uncertainty, freshness and operational limitations; prevent unreviewed output from becoming a command, authoritative record or safety conclusion.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Output evaluation; safety/content test record; uncertainty-display evidence; exception and escalation record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is pii leakage prevention scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D5-CTL-03 - COPYRIGHT DETECTION

Field	Energy-sector treatment
Domain	D5: CONTENT SAFETY & OUTPUT INTEGRITY
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D5-CTL-03
Interpretation	In an energy environment, copyright detection is implemented by applying test outputs for unsafe, deceptive, confidential or operationally misleading content; communicate uncertainty and freshness. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Display confidence, uncertainty, freshness and operational limitations; prevent unreviewed output from becoming a command, authoritative record or safety conclusion.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and

Field	Energy-sector treatment
	information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Output evaluation; safety/content test record; uncertainty-display evidence; exception and escalation record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is copyright detection scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D5-CTL-04 - AI WATERMARKING ROBUSTNESS

Field	Energy-sector treatment
Domain	D5: CONTENT SAFETY & OUTPUT INTEGRITY
Implementation sequence	A - Advanced / contextual
Normative reference	GAISF-NOR-001, D5-CTL-04
Interpretation	In an energy environment, ai watermarking robustness is implemented by applying test outputs for unsafe, deceptive, confidential or operationally misleading content; communicate uncertainty and freshness. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Display confidence, uncertainty, freshness and operational limitations; prevent unreviewed output from becoming a command, authoritative record or safety conclusion.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Output evaluation; safety/content test record; uncertainty-display evidence; exception and escalation record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is ai watermarking robustness scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D5-CTL-05 - PRIVACY-BY-DESIGN VERIFICATION

Field	Energy-sector treatment
Domain	D5: CONTENT SAFETY & OUTPUT INTEGRITY

Field	Energy-sector treatment
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D5-CTL-05
Interpretation	In an energy environment, privacy-by-design verification is implemented by applying test outputs for unsafe, deceptive, confidential or operationally misleading content; communicate uncertainty and freshness. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Display confidence, uncertainty, freshness and operational limitations; prevent unreviewed output from becoming a command, authoritative record or safety conclusion.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Output evaluation; safety/content test record; uncertainty-display evidence; exception and escalation record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is privacy-by-design verification scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D5-CTL-06 - PRIVACY-PRESERVING ML VALIDATION

Field	Energy-sector treatment
Domain	D5: CONTENT SAFETY & OUTPUT INTEGRITY
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D5-CTL-06
Interpretation	In an energy environment, privacy-preserving ml validation is implemented by applying test outputs for unsafe, deceptive, confidential or operationally misleading content; communicate uncertainty and freshness. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Display confidence, uncertainty, freshness and operational limitations; prevent unreviewed output from becoming a command, authoritative record or safety conclusion.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability

Field	Energy-sector treatment
	function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Output evaluation; safety/content test record; uncertainty-display evidence; exception and escalation record. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is privacy-preserving ml validation scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D6-CTL-01 - HUMAN-IN-THE-LOOP FOR HIGH-RISK ACTIONS

Field	Energy-sector treatment
Domain	D6: GOVERNANCE, ACCOUNTABILITY & HUMAN OVERSIGHT
Implementation sequence	F - Foundational
Normative reference	GAISF-NOR-001, D6-CTL-01
Interpretation	In an energy environment, human-in-the-loop for high-risk actions is implemented by applying assign accountable owners, decision rights, audit trails, incident management, rollback and lifecycle governance. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use cross-functional gates involving operations, OT security, safety/reliability, AI governance and accountable ownership; retain decisions, exceptions and review triggers.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Approved policy or procedure; named accountability; audit trail; incident, rollback, change and governance decision records. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is human-in-the-loop for high-risk actions scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D6-CTL-02 - AUDIT TRAIL COMPLETENESS

Field	Energy-sector treatment
Domain	D6: GOVERNANCE, ACCOUNTABILITY & HUMAN OVERSIGHT
Implementation sequence	F - Foundational
Normative reference	GAISSF-NOR-001, D6-CTL-02
Interpretation	In an energy environment, audit trail completeness is implemented by applying assign accountable owners, decision rights, audit trails, incident management, rollback and lifecycle governance. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use cross-functional gates involving operations, OT security, safety/reliability, AI governance and accountable ownership; retain decisions, exceptions and review triggers.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Approved policy or procedure; named accountability; audit trail; incident, rollback, change and governance decision records. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is audit trail completeness scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D6-CTL-03 - AI MODEL CARD COMPLETENESS

Field	Energy-sector treatment
Domain	D6: GOVERNANCE, ACCOUNTABILITY & HUMAN OVERSIGHT
Implementation sequence	F - Foundational
Normative reference	GAISSF-NOR-001, D6-CTL-03
Interpretation	In an energy environment, ai model card completeness is implemented by applying assign accountable owners, decision rights, audit trails, incident management, rollback and lifecycle governance. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use cross-functional gates involving operations, OT security, safety/reliability, AI governance and accountable ownership; retain decisions, exceptions and review triggers.

Field	Energy-sector treatment
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Approved policy or procedure; named accountability; audit trail; incident, rollback, change and governance decision records. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is ai model card completeness scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D6-CTL-04 - AI INCIDENT RESPONSE READINESS

Field	Energy-sector treatment
Domain	D6: GOVERNANCE, ACCOUNTABILITY & HUMAN OVERSIGHT
Implementation sequence	F - Foundational
Normative reference	GAISSE-NOR-001, D6-CTL-04
Interpretation	In an energy environment, ai incident response readiness is implemented by applying assign accountable owners, decision rights, audit trails, incident management, rollback and lifecycle governance. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use cross-functional gates involving operations, OT security, safety/reliability, AI governance and accountable ownership; retain decisions, exceptions and review triggers.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Approved policy or procedure; named accountability; audit trail; incident, rollback, change and governance decision records. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is ai incident response readiness scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created

Field	Energy-sector treatment
	only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D6-CTL-05 - MODEL DEPRECATION & DECOMMISSIONING

Field	Energy-sector treatment
Domain	D6: GOVERNANCE, ACCOUNTABILITY & HUMAN OVERSIGHT
Implementation sequence	F - Foundational
Normative reference	GAISSF-NOR-001, D6-CTL-05
Interpretation	In an energy environment, model deprecation & decommissioning is implemented by applying assign accountable owners, decision rights, audit trails, incident management, rollback and lifecycle governance. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use cross-functional gates involving operations, OT security, safety/reliability, AI governance and accountable ownership; retain decisions, exceptions and review triggers.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Approved policy or procedure; named accountability; audit trail; incident, rollback, change and governance decision records. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is model deprecation & decommissioning scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D6-CTL-06 - THIRD-PARTY AI VENDOR GOVERNANCE

Field	Energy-sector treatment
Domain	D6: GOVERNANCE, ACCOUNTABILITY & HUMAN OVERSIGHT
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D6-CTL-06
Interpretation	In an energy environment, third-party ai vendor governance is implemented by applying assign accountable owners, decision rights, audit trails, incident management, rollback and lifecycle governance. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty

Field	Energy-sector treatment
	alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use cross-functional gates involving operations, OT security, safety/reliability, AI governance and accountable ownership; retain decisions, exceptions and review triggers.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Approved policy or procedure; named accountability; audit trail; incident, rollback, change and governance decision records. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is third-party ai vendor governance scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D6-CTL-07 - AI RESILIENCE & BUSINESS CONTINUITY

Field	Energy-sector treatment
Domain	D6: GOVERNANCE, ACCOUNTABILITY & HUMAN OVERSIGHT
Implementation sequence	O - Operational
Normative reference	GAISSE-NOR-001, D6-CTL-07
Interpretation	In an energy environment, ai resilience & business continuity is implemented by applying assign accountable owners, decision rights, audit trails, incident management, rollback and lifecycle governance. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Use cross-functional gates involving operations, OT security, safety/reliability, AI governance and accountable ownership; retain decisions, exceptions and review triggers.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Approved policy or procedure; named accountability; audit trail; incident, rollback, change and governance decision records. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate;

Field	Energy-sector treatment
	rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is ai resilience & business continuity scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D7-CTL-H01 - AI-GENERATED PHISHING SIMULATION

Field	Energy-sector treatment
Domain	D7: HUMAN & SOCIETAL HARMS
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D7-CTL-H01
Interpretation	In an energy environment, ai-generated phishing simulation is implemented by applying design competent oversight, out-of-band verification, escalation, workload controls and resistance to impersonation or automation bias. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Train through realistic exercises; test override, dissent, escalation and out-of-band verification; monitor automation bias and workload.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Human-factors assessment; training and exercise record; out-of-band verification evidence; override and escalation records. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is ai-generated phishing simulation scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D7-CTL-H02 - DEEPFAKE DETECTION TRAINING

Field	Energy-sector treatment
Domain	D7: HUMAN & SOCIETAL HARMS
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D7-CTL-H02
Interpretation	In an energy environment, deepfake detection training is implemented by applying design competent oversight, out-of-band verification, escalation, workload controls and resistance to impersonation or automation bias. Implementation intensity is determined by operational

Field	Energy-sector treatment
	authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Train through realistic exercises; test override, dissent, escalation and out-of-band verification; monitor automation bias and workload.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Human-factors assessment; training and exercise record; out-of-band verification evidence; override and escalation records. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is deepfake detection training scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D7-CTL-H03 - OUT-OF-BAND AUTHENTICATION

Field	Energy-sector treatment
Domain	D7: HUMAN & SOCIETAL HARMS
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D7-CTL-H03
Interpretation	In an energy environment, out-of-band authentication is implemented by applying design competent oversight, out-of-band verification, escalation, workload controls and resistance to impersonation or automation bias. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Train through realistic exercises; test override, dissent, escalation and out-of-band verification; monitor automation bias and workload.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Human-factors assessment; training and exercise record; out-of-band verification evidence; override and escalation records. Retention follows applicable law, contract, safety case and organisational records policy.

Field	Energy-sector treatment
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is out-of-band authentication scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D7-CTL-H04 - AI SOCIAL ENGINEERING IR

Field	Energy-sector treatment
Domain	D7: HUMAN & SOCIETAL HARMS
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D7-CTL-H04
Interpretation	In an energy environment, ai social engineering ir is implemented by applying design competent oversight, out-of-band verification, escalation, workload controls and resistance to impersonation or automation bias. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Train through realistic exercises; test override, dissent, escalation and out-of-band verification; monitor automation bias and workload.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Human-factors assessment; training and exercise record; out-of-band verification evidence; override and escalation records. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is ai social engineering ir scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D7-CTL-H05 - AI-ENHANCED EXTERNAL ATTACK DEFENSE

Field	Energy-sector treatment
Domain	D7: HUMAN & SOCIETAL HARMS
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D7-CTL-H05
Interpretation	In an energy environment, ai-enhanced external attack defense is implemented by applying design competent oversight, out-of-band verification, escalation, workload controls and resistance to

Field	Energy-sector treatment
	impersonation or automation bias. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Train through realistic exercises; test override, dissent, escalation and out-of-band verification; monitor automation bias and workload.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Human-factors assessment; training and exercise record; out-of-band verification evidence; override and escalation records. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is ai-enhanced external attack defense scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D8-CTL-01 - EU AI ACT RISK TIER MAPPING

Field	Energy-sector treatment
Domain	D8: REGULATORY ALIGNMENT & COMPLIANCE
Implementation sequence	F - Foundational
Normative reference	GAISF-NOR-001, D8-CTL-01
Interpretation	In an energy environment, eu ai act risk tier mapping is implemented by applying maintain applicable-obligation records, evidence, assessments and incident reporting without asserting cross-framework equivalence. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Maintain a jurisdiction-specific obligations register; use crosswalks only as informative support; obtain qualified interpretation for legal or regulatory claims.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Obligations register; applicability rationale; mapping record; assessment

Field	Energy-sector treatment
	report; incident/reporting evidence where legally applicable. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is eu ai act risk tier mapping scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D8-CTL-02 - ISO 42001 GAP ANALYSIS

Field	Energy-sector treatment
Domain	D8: REGULATORY ALIGNMENT & COMPLIANCE
Implementation sequence	F - Foundational
Normative reference	GAISSF-NOR-001, D8-CTL-02
Interpretation	In an energy environment, iso 42001 gap analysis is implemented by applying maintain applicable-obligation records, evidence, assessments and incident reporting without asserting cross-framework equivalence. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Maintain a jurisdiction-specific obligations register; use crosswalks only as informative support; obtain qualified interpretation for legal or regulatory claims.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Obligations register; applicability rationale; mapping record; assessment report; incident/reporting evidence where legally applicable. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is iso 42001 gap analysis scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D8-CTL-03 - GPAI TECHNICAL DOCUMENTATION VERIFICATION

Field	Energy-sector treatment
Domain	D8: REGULATORY ALIGNMENT & COMPLIANCE

Field	Energy-sector treatment
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D8-CTL-03
Interpretation	In an energy environment, gpai technical documentation verification is implemented by applying maintain applicable-obligation records, evidence, assessments and incident reporting without asserting cross-framework equivalence. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Maintain a jurisdiction-specific obligations register; use crosswalks only as informative support; obtain qualified interpretation for legal or regulatory claims.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Obligations register; applicability rationale; mapping record; assessment report; incident/reporting evidence where legally applicable. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is gpai technical documentation verification scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D8-CTL-04 - DORA ICT INCIDENT REPORTING (FINANCIAL SECTOR)

Field	Energy-sector treatment
Domain	D8: REGULATORY ALIGNMENT & COMPLIANCE
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D8-CTL-04
Interpretation	In an energy environment, dora ict incident reporting (financial sector) is implemented by applying maintain applicable-obligation records, evidence, assessments and incident reporting without asserting cross-framework equivalence. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Maintain a jurisdiction-specific obligations register; use crosswalks only as informative support; obtain qualified interpretation for legal or regulatory claims.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.

Field	Energy-sector treatment
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Obligations register; applicability rationale; mapping record; assessment report; incident/reporting evidence where legally applicable. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is dora ict incident reporting (financial sector) scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D8-CTL-05 - NIST SP 800-218A COMPLIANCE CHECK

Field	Energy-sector treatment
Domain	D8: REGULATORY ALIGNMENT & COMPLIANCE
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D8-CTL-05
Interpretation	In an energy environment, nist sp 800-218a compliance check is implemented by applying maintain applicable-obligation records, evidence, assessments and incident reporting without asserting cross-framework equivalence. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Maintain a jurisdiction-specific obligations register; use crosswalks only as informative support; obtain qualified interpretation for legal or regulatory claims.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Obligations register; applicability rationale; mapping record; assessment report; incident/reporting evidence where legally applicable. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is nist sp 800-218a compliance check scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.

Field	Energy-sector treatment
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D9-CTL-01 - PHYSICAL HARM BOUNDARY ENFORCEMENT

Field	Energy-sector treatment
Domain	D9: PHYSICAL AI SAFETY
Implementation sequence	F - Foundational
Normative reference	GAISSF-NOR-001, D9-CTL-01
Interpretation	In an energy environment, physical harm boundary enforcement is implemented by applying preserve independent protection layers, safe states, sensor integrity, command verification, override, fail-safe and forensic evidence. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Preserve independent interlocks and protection layers; use safe states, manual override, command verification, sensor plausibility, deterministic fallback and tested restoration.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Hazard and boundary analysis; independent-protection evidence; safe-state, override, sensor, command and recovery test reports. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is physical harm boundary enforcement scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D9-CTL-02 - SAFE STATE AND GRACEFUL DEGRADATION

Field	Energy-sector treatment
Domain	D9: PHYSICAL AI SAFETY
Implementation sequence	F - Foundational
Normative reference	GAISSF-NOR-001, D9-CTL-02
Interpretation	In an energy environment, safe state and graceful degradation is implemented by applying preserve independent protection layers, safe states, sensor integrity, command verification, override, fail-safe and forensic evidence. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative

Field	Energy-sector treatment
	testing; operational monitoring; reconstructable evidence.
Recommended practices	Preserve independent interlocks and protection layers; use safe states, manual override, command verification, sensor plausibility, deterministic fallback and tested restoration.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Hazard and boundary analysis; independent-protection evidence; safe-state, override, sensor, command and recovery test reports. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is safe state and graceful degradation scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D9-CTL-03 - HUMAN OVERRIDE AND EMERGENCY STOP

Field	Energy-sector treatment
Domain	D9: PHYSICAL AI SAFETY
Implementation sequence	F - Foundational
Normative reference	GAISSF-NOR-001, D9-CTL-03
Interpretation	In an energy environment, human override and emergency stop is implemented by applying preserve independent protection layers, safe states, sensor integrity, command verification, override, fail-safe and forensic evidence. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Preserve independent interlocks and protection layers; use safe states, manual override, command verification, sensor plausibility, deterministic fallback and tested restoration.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Hazard and boundary analysis; independent-protection evidence; safe-state, override, sensor, command and recovery test reports. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is human override and emergency stop scoped, implemented,

Field	Energy-sector treatment
	tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D9-CTL-04 - CYBER-PHYSICAL ATTACK DETECTION

Field	Energy-sector treatment
Domain	D9: PHYSICAL AI SAFETY
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D9-CTL-04
Interpretation	In an energy environment, cyber-physical attack detection is implemented by applying preserve independent protection layers, safe states, sensor integrity, command verification, override, fail-safe and forensic evidence. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Preserve independent interlocks and protection layers; use safe states, manual override, command verification, sensor plausibility, deterministic fallback and tested restoration.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Hazard and boundary analysis; independent-protection evidence; safe-state, override, sensor, command and recovery test reports. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is cyber-physical attack detection scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D9-CTL-05 - PHYSICAL ENVIRONMENT INTEGRITY MONITORING

Field	Energy-sector treatment
Domain	D9: PHYSICAL AI SAFETY
Implementation sequence	O - Operational
Normative reference	GAISSF-NOR-001, D9-CTL-05
Interpretation	In an energy environment, physical environment integrity monitoring is implemented by applying preserve independent protection layers, safe states, sensor integrity, command verification, override, fail-safe and forensic evidence. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality,

Field	Energy-sector treatment
	connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Preserve independent interlocks and protection layers; use safe states, manual override, command verification, sensor plausibility, deterministic fallback and tested restoration.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Hazard and boundary analysis; independent-protection evidence; safe-state, override, sensor, command and recovery test reports. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is physical environment integrity monitoring scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D9-CTL-06 - ACTUATOR COMMAND VERIFICATION

Field	Energy-sector treatment
Domain	D9: PHYSICAL AI SAFETY
Implementation sequence	A - Advanced / contextual
Normative reference	GAISSF-NOR-001, D9-CTL-06
Interpretation	In an energy environment, actuator command verification is implemented by applying preserve independent protection layers, safe states, sensor integrity, command verification, override, fail-safe and forensic evidence. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Preserve independent interlocks and protection layers; use safe states, manual override, command verification, sensor plausibility, deterministic fallback and tested restoration.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Hazard and boundary analysis; independent-protection evidence; safe-

Field	Energy-sector treatment
	state, override, sensor, command and recovery test reports. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is actuator command verification scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

D9-CTL-07 - PHYSICAL INCIDENT EVIDENCE PRESERVATION

Field	Energy-sector treatment
Domain	D9: PHYSICAL AI SAFETY
Implementation sequence	A - Advanced / contextual
Normative reference	GAISSF-NOR-001, D9-CTL-07
Interpretation	In an energy environment, physical incident evidence preservation is implemented by applying preserve independent protection layers, safe states, sensor integrity, command verification, override, fail-safe and forensic evidence. Implementation intensity is determined by operational authority, safety and reliability consequence, criticality, connectivity, reversibility and dependency concentration.
Applicability / N/A rule	Assess for every in-scope AI system. Marking Not Applicable requires a defined system boundary, evidence that the addressed risk is absent, an accountable approval, and a reassessment trigger. Technical difficulty alone is not a valid N/A rationale.
Required outcomes	Documented applicability; accountable owner; approved implementation or time-bound compensating measure; representative testing; operational monitoring; reconstructable evidence.
Recommended practices	Preserve independent interlocks and protection layers; use safe states, manual override, command verification, sensor plausibility, deterministic fallback and tested restoration.
OT/ICS considerations	Protect deterministic control and availability; avoid unplanned operational disruption; respect maintenance windows, legacy protocols, segmentation, latency and offline operation.
Safety/reliability	AI must not bypass an independent safety, protection or reliability function. Test degraded modes, communications loss, stale data, fallback, rollback and recovery.
Human oversight	Define the human decision point, competence, available time and information, approval authority, override mechanism, escalation threshold and evidence of action.
Evidence requirements	Hazard and boundary analysis; independent-protection evidence; safe-state, override, sensor, command and recovery test reports. Retention follows applicable law, contract, safety case and organisational records policy.
Metrics	Applicability coverage; test completion and pass rate; rejected/overridden actions; drift or integrity alerts; open findings; fallback invocations; incidents and near misses.
Assessment question	How is physical incident evidence preservation scoped, implemented, tested, monitored and evidenced for this system's energy-sector consequence and authority?
Common failure modes	Policy-only evidence; copied supplier assertions; stale applicability; missing OT or safety constraints; untested fallback; evidence created only for assessment.
Compensating measures	Document equivalent risk reduction, accountable approval, monitoring, expiry and a remediation or reassessment date.
Limitations	This interpretation is informative and requires system-, engineering- and jurisdiction-specific review.

Energy AI use cases and differentiated mappings

Mappings are informative implementation relationships. They do not narrow the full GAISSF applicability assessment.

UC-01 - Short-term load forecasting

Field	Content
Subsector	Electricity
Authority Level	Operational decision support
Objective	Use AI to support short-term load forecasting while preserving accountable, bounded and recoverable operation.
Data Inputs	Telemetry, asset, operational, environmental, market, customer or security data as applicable; lineage and freshness are recorded.
Human Decision Point	Named operator or accountable owner reviews or authorises material action according to authority class.
Failure Modes	Incorrect prediction; stale, missing or manipulated data; distribution shift; dependency loss; automation bias; unauthorised or unsafe action.
Applicable Controls	D1-CTL-01; D1-CTL-03; D1-CTL-04; D4-CTL-02; D5-CTL-01; D6-CTL-01; D6-CTL-02; D6-CTL-04; D6-CTL-05; D8-CTL-01
Mapping Rationale	Selected for the operational decision support authority class, electricity context and the stated data, human-oversight, dependency and consequence profile.
Evidence	System risk assessment; architecture and authority record; representative validation; approval; monitoring; override, fallback and incident evidence.
Fallback	Revert to an approved deterministic procedure, trusted prior version or manual operation without bypassing independent protection.
Residual Risk	Accepted by the accountable owner after security, operations and applicable safety/reliability review.
Notably Absent	No claim that the use of AI improves reliability, safety, fairness or efficiency without system-specific operational evidence.

UC-02 - Renewable generation forecasting

Field	Content
Subsector	Electricity
Authority Level	Operational advisory
Objective	Use AI to support renewable generation forecasting while preserving accountable, bounded and recoverable operation.
Data Inputs	Telemetry, asset, operational, environmental, market, customer or security data as applicable; lineage and freshness are recorded.
Human Decision Point	Named operator or accountable owner reviews or authorises material action according to authority class.
Failure Modes	Incorrect prediction; stale, missing or manipulated data; distribution shift; dependency loss; automation bias; unauthorised or unsafe action.
Applicable Controls	D1-CTL-01; D1-CTL-03; D1-CTL-04; D4-CTL-02; D5-CTL-01; D6-CTL-02; D6-CTL-04; D8-CTL-01
Mapping Rationale	Selected for the operational advisory authority class, electricity context and the stated data, human-oversight, dependency and consequence profile.
Evidence	System risk assessment; architecture and authority record; representative validation; approval; monitoring; override, fallback and incident evidence.
Fallback	Revert to an approved deterministic procedure, trusted prior version or manual operation without bypassing independent protection.
Residual Risk	Accepted by the accountable owner after security, operations and applicable safety/reliability review.

Field	Content
Notably Absent	No claim that the use of AI improves reliability, safety, fairness or efficiency without system-specific operational evidence.

UC-03 - Transformer health analytics

Field	Content
Subsector	Electricity
Authority Level	Predictive maintenance
Objective	Use AI to support transformer health analytics while preserving accountable, bounded and recoverable operation.
Data Inputs	Telemetry, asset, operational, environmental, market, customer or security data as applicable; lineage and freshness are recorded.
Human Decision Point	Named operator or accountable owner reviews or authorises material action according to authority class.
Failure Modes	Incorrect prediction; stale, missing or manipulated data; distribution shift; dependency loss; automation bias; unauthorised or unsafe action.
Applicable Controls	D1-CTL-01; D1-CTL-03; D1-CTL-08; D4-CTL-01; D4-CTL-02; D6-CTL-02; D6-CTL-04; D6-CTL-05; D9-CTL-03
Mapping Rationale	Selected for the predictive maintenance authority class, electricity context and the stated data, human-oversight, dependency and consequence profile.
Evidence	System risk assessment; architecture and authority record; representative validation; approval; monitoring; override, fallback and incident evidence.
Fallback	Revert to an approved deterministic procedure, trusted prior version or manual operation without bypassing independent protection.
Residual Risk	Accepted by the accountable owner after security, operations and applicable safety/reliability review.
Notably Absent	No claim that the use of AI improves reliability, safety, fairness or efficiency without system-specific operational evidence.

UC-04 - Fault detection and localisation

Field	Content
Subsector	Electricity
Authority Level	High-impact advisory
Objective	Use AI to support fault detection and localisation while preserving accountable, bounded and recoverable operation.
Data Inputs	Telemetry, asset, operational, environmental, market, customer or security data as applicable; lineage and freshness are recorded.
Human Decision Point	Named operator or accountable owner reviews or authorises material action according to authority class.
Failure Modes	Incorrect prediction; stale, missing or manipulated data; distribution shift; dependency loss; automation bias; unauthorised or unsafe action.
Applicable Controls	D1-CTL-01; D1-CTL-03; D2-CTL-04; D3-CTL-01; D5-CTL-01; D6-CTL-01; D6-CTL-02; D6-CTL-04; D9-CTL-02; D9-CTL-03
Mapping Rationale	Selected for the high-impact advisory authority class, electricity context and the stated data, human-oversight, dependency and consequence profile.
Evidence	System risk assessment; architecture and authority record; representative validation; approval; monitoring; override, fallback and incident evidence.
Fallback	Revert to an approved deterministic procedure, trusted prior version or manual operation without bypassing independent protection.
Residual Risk	Accepted by the accountable owner after security, operations and applicable safety/reliability review.

Field	Content
Notably Absent	No claim that the use of AI improves reliability, safety, fairness or efficiency without system-specific operational evidence.

UC-05 - Outage prediction

Field	Content
Subsector	Electricity
Authority Level	Operational advisory
Objective	Use AI to support outage prediction while preserving accountable, bounded and recoverable operation.
Data Inputs	Telemetry, asset, operational, environmental, market, customer or security data as applicable; lineage and freshness are recorded.
Human Decision Point	Named operator or accountable owner reviews or authorises material action according to authority class.
Failure Modes	Incorrect prediction; stale, missing or manipulated data; distribution shift; dependency loss; automation bias; unauthorised or unsafe action.
Applicable Controls	D1-CTL-01; D1-CTL-03; D4-CTL-02; D5-CTL-01; D6-CTL-02; D6-CTL-04; D7-CTL-04; D8-CTL-01
Mapping Rationale	Selected for the operational advisory authority class, electricity context and the stated data, human-oversight, dependency and consequence profile.
Evidence	System risk assessment; architecture and authority record; representative validation; approval; monitoring; override, fallback and incident evidence.
Fallback	Revert to an approved deterministic procedure, trusted prior version or manual operation without bypassing independent protection.
Residual Risk	Accepted by the accountable owner after security, operations and applicable safety/reliability review.
Notably Absent	No claim that the use of AI improves reliability, safety, fairness or efficiency without system-specific operational evidence.

UC-06 - Restoration prioritisation

Field	Content
Subsector	Electricity
Authority Level	High-impact decision support
Objective	Use AI to support restoration prioritisation while preserving accountable, bounded and recoverable operation.
Data Inputs	Telemetry, asset, operational, environmental, market, customer or security data as applicable; lineage and freshness are recorded.
Human Decision Point	Named operator or accountable owner reviews or authorises material action according to authority class.
Failure Modes	Incorrect prediction; stale, missing or manipulated data; distribution shift; dependency loss; automation bias; unauthorised or unsafe action.
Applicable Controls	D1-CTL-01; D1-CTL-03; D3-CTL-01; D5-CTL-01; D5-CTL-02; D6-CTL-01; D6-CTL-02; D6-CTL-04; D7-CTL-04; D8-CTL-01; D9-CTL-02
Mapping Rationale	Selected for the high-impact decision support authority class, electricity context and the stated data, human-oversight, dependency and consequence profile.
Evidence	System risk assessment; architecture and authority record; representative validation; approval; monitoring; override, fallback and incident evidence.
Fallback	Revert to an approved deterministic procedure, trusted prior version or manual operation without bypassing independent protection.
Residual Risk	Accepted by the accountable owner after security, operations and applicable safety/reliability review.

Field	Content
Notably Absent	No claim that the use of AI improves reliability, safety, fairness or efficiency without system-specific operational evidence.

UC-07 - Voltage and reactive-power optimisation

Field	Content
Subsector	Electricity
Authority Level	High-impact automated
Objective	Use AI to support voltage and reactive-power optimisation while preserving accountable, bounded and recoverable operation.
Data Inputs	Telemetry, asset, operational, environmental, market, customer or security data as applicable; lineage and freshness are recorded.
Human Decision Point	Named operator or accountable owner reviews or authorises material action according to authority class.
Failure Modes	Incorrect prediction; stale, missing or manipulated data; distribution shift; dependency loss; automation bias; unauthorised or unsafe action.
Applicable Controls	D1-CTL-01; D1-CTL-03; D3-CTL-01; D3-CTL-02; D3-CTL-04; D6-CTL-01; D6-CTL-02; D6-CTL-04; D6-CTL-05; D9-CTL-01; D9-CTL-02; D9-CTL-03; D9-CTL-06
Mapping Rationale	Selected for the high-impact automated authority class, electricity context and the stated data, human-oversight, dependency and consequence profile.
Evidence	System risk assessment; architecture and authority record; representative validation; approval; monitoring; override, fallback and incident evidence.
Fallback	Revert to an approved deterministic procedure, trusted prior version or manual operation without bypassing independent protection.
Residual Risk	Accepted by the accountable owner after security, operations and applicable safety/reliability review.
Notably Absent	No claim that the use of AI improves reliability, safety, fairness or efficiency without system-specific operational evidence.

UC-08 - Distributed energy-resource orchestration

Field	Content
Subsector	Electricity
Authority Level	Closed-loop operational
Objective	Use AI to support distributed energy-resource orchestration while preserving accountable, bounded and recoverable operation.
Data Inputs	Telemetry, asset, operational, environmental, market, customer or security data as applicable; lineage and freshness are recorded.
Human Decision Point	Named operator or accountable owner reviews or authorises material action according to authority class.
Failure Modes	Incorrect prediction; stale, missing or manipulated data; distribution shift; dependency loss; automation bias; unauthorised or unsafe action.
Applicable Controls	D1-CTL-01; D1-CTL-03; D2-CTL-04; D3-CTL-01; D3-CTL-02; D3-CTL-04; D4-CTL-01; D4-CTL-02; D6-CTL-01; D6-CTL-02; D6-CTL-04; D6-CTL-05; D9-CTL-01; D9-CTL-02; D9-CTL-03; D9-CTL-04; D9-CTL-05; D9-CTL-06
Mapping Rationale	Selected for the closed-loop operational authority class, electricity context and the stated data, human-oversight, dependency and consequence profile.
Evidence	System risk assessment; architecture and authority record; representative validation; approval; monitoring; override, fallback and incident evidence.

Field	Content
Fallback	Revert to an approved deterministic procedure, trusted prior version or manual operation without bypassing independent protection.
Residual Risk	Accepted by the accountable owner after security, operations and applicable safety/reliability review.
Notably Absent	No claim that the use of AI improves reliability, safety, fairness or efficiency without system-specific operational evidence.

UC-09 - Battery energy-storage optimisation

Field	Content
Subsector	Storage
Authority Level	High-impact automated
Objective	Use AI to support battery energy-storage optimisation while preserving accountable, bounded and recoverable operation.
Data Inputs	Telemetry, asset, operational, environmental, market, customer or security data as applicable; lineage and freshness are recorded.
Human Decision Point	Named operator or accountable owner reviews or authorises material action according to authority class.
Failure Modes	Incorrect prediction; stale, missing or manipulated data; distribution shift; dependency loss; automation bias; unauthorised or unsafe action.
Applicable Controls	D1-CTL-01; D1-CTL-03; D3-CTL-01; D3-CTL-02; D6-CTL-01; D6-CTL-02; D6-CTL-04; D6-CTL-05; D9-CTL-01; D9-CTL-02; D9-CTL-03; D9-CTL-06
Mapping Rationale	Selected for the high-impact automated authority class, storage context and the stated data, human-oversight, dependency and consequence profile.
Evidence	System risk assessment; architecture and authority record; representative validation; approval; monitoring; override, fallback and incident evidence.
Fallback	Revert to an approved deterministic procedure, trusted prior version or manual operation without bypassing independent protection.
Residual Risk	Accepted by the accountable owner after security, operations and applicable safety/reliability review.
Notably Absent	No claim that the use of AI improves reliability, safety, fairness or efficiency without system-specific operational evidence.

UC-10 - Pipeline anomaly detection

Field	Content
Subsector	Oil and gas
Authority Level	Safety-relevant advisory
Objective	Use AI to support pipeline anomaly detection while preserving accountable, bounded and recoverable operation.
Data Inputs	Telemetry, asset, operational, environmental, market, customer or security data as applicable; lineage and freshness are recorded.
Human Decision Point	Named operator or accountable owner reviews or authorises material action according to authority class.
Failure Modes	Incorrect prediction; stale, missing or manipulated data; distribution shift; dependency loss; automation bias; unauthorised or unsafe action.
Applicable Controls	D1-CTL-01; D1-CTL-03; D1-CTL-08; D2-CTL-04; D4-CTL-02; D6-CTL-02; D6-CTL-04; D8-CTL-01; D9-CTL-02; D9-CTL-03
Mapping Rationale	Selected for the safety-relevant advisory authority class, oil and gas context and the stated data, human-oversight, dependency and consequence profile.
Evidence	System risk assessment; architecture and authority record; representative validation; approval; monitoring; override,

Field	Content
	fallback and incident evidence.
Fallback	Revert to an approved deterministic procedure, trusted prior version or manual operation without bypassing independent protection.
Residual Risk	Accepted by the accountable owner after security, operations and applicable safety/reliability review.
Notably Absent	No claim that the use of AI improves reliability, safety, fairness or efficiency without system-specific operational evidence.

UC-11 - Leak-detection support

Field	Content
Subsector	Oil and gas
Authority Level	Safety-relevant decision support
Objective	Use AI to support leak-detection support while preserving accountable, bounded and recoverable operation.
Data Inputs	Telemetry, asset, operational, environmental, market, customer or security data as applicable; lineage and freshness are recorded.
Human Decision Point	Named operator or accountable owner reviews or authorises material action according to authority class.
Failure Modes	Incorrect prediction; stale, missing or manipulated data; distribution shift; dependency loss; automation bias; unauthorised or unsafe action.
Applicable Controls	D1-CTL-01; D1-CTL-03; D3-CTL-01; D5-CTL-01; D6-CTL-01; D6-CTL-02; D6-CTL-04; D7-CTL-04; D9-CTL-01; D9-CTL-02; D9-CTL-03
Mapping Rationale	Selected for the safety-relevant decision support authority class, oil and gas context and the stated data, human-oversight, dependency and consequence profile.
Evidence	System risk assessment; architecture and authority record; representative validation; approval; monitoring; override, fallback and incident evidence.
Fallback	Revert to an approved deterministic procedure, trusted prior version or manual operation without bypassing independent protection.
Residual Risk	Accepted by the accountable owner after security, operations and applicable safety/reliability review.
Notably Absent	No claim that the use of AI improves reliability, safety, fairness or efficiency without system-specific operational evidence.

UC-12 - Refinery process optimisation

Field	Content
Subsector	Oil and gas
Authority Level	High-impact automated
Objective	Use AI to support refinery process optimisation while preserving accountable, bounded and recoverable operation.
Data Inputs	Telemetry, asset, operational, environmental, market, customer or security data as applicable; lineage and freshness are recorded.
Human Decision Point	Named operator or accountable owner reviews or authorises material action according to authority class.
Failure Modes	Incorrect prediction; stale, missing or manipulated data; distribution shift; dependency loss; automation bias; unauthorised or unsafe action.
Applicable Controls	D1-CTL-01; D1-CTL-03; D3-CTL-01; D3-CTL-02; D4-CTL-01; D6-CTL-01; D6-CTL-02; D6-CTL-04; D6-CTL-05; D9-CTL-01; D9-CTL-02; D9-CTL-03; D9-CTL-06
Mapping Rationale	Selected for the high-impact automated authority class, oil and gas context and the stated data, human-oversight, dependency and consequence profile.
Evidence	System risk assessment; architecture and authority record;

Field	Content
	representative validation; approval; monitoring; override, fallback and incident evidence.
Fallback	Revert to an approved deterministic procedure, trusted prior version or manual operation without bypassing independent protection.
Residual Risk	Accepted by the accountable owner after security, operations and applicable safety/reliability review.
Notably Absent	No claim that the use of AI improves reliability, safety, fairness or efficiency without system-specific operational evidence.

UC-13 - Energy trading and bidding

Field	Content
Subsector	Markets
Authority Level	Automated financial decision
Objective	Use AI to support energy trading and bidding while preserving accountable, bounded and recoverable operation.
Data Inputs	Telemetry, asset, operational, environmental, market, customer or security data as applicable; lineage and freshness are recorded.
Human Decision Point	Named operator or accountable owner reviews or authorises material action according to authority class.
Failure Modes	Incorrect prediction; stale, missing or manipulated data; distribution shift; dependency loss; automation bias; unauthorised or unsafe action.
Applicable Controls	D1-CTL-01; D1-CTL-03; D2-CTL-01; D3-CTL-01; D3-CTL-02; D5-CTL-01; D5-CTL-02; D5-CTL-04; D6-CTL-01; D6-CTL-02; D6-CTL-04; D7-CTL-05; D8-CTL-01; D8-CTL-02; D8-CTL-04
Mapping Rationale	Selected for the automated financial decision authority class, markets context and the stated data, human-oversight, dependency and consequence profile.
Evidence	System risk assessment; architecture and authority record; representative validation; approval; monitoring; override, fallback and incident evidence.
Fallback	Revert to an approved deterministic procedure, trusted prior version or manual operation without bypassing independent protection.
Residual Risk	Accepted by the accountable owner after security, operations and applicable safety/reliability review.
Notably Absent	No claim that the use of AI improves reliability, safety, fairness or efficiency without system-specific operational evidence.

UC-14 - Demand-response targeting

Field	Content
Subsector	Retail electricity
Authority Level	Customer and operational decision
Objective	Use AI to support demand-response targeting while preserving accountable, bounded and recoverable operation.
Data Inputs	Telemetry, asset, operational, environmental, market, customer or security data as applicable; lineage and freshness are recorded.
Human Decision Point	Named operator or accountable owner reviews or authorises material action according to authority class.
Failure Modes	Incorrect prediction; stale, missing or manipulated data; distribution shift; dependency loss; automation bias; unauthorised or unsafe action.
Applicable Controls	D1-CTL-01; D1-CTL-03; D5-CTL-01; D5-CTL-02; D5-CTL-05; D6-CTL-01; D6-CTL-02; D6-CTL-04; D7-CTL-04; D8-CTL-01; D8-CTL-02
Mapping Rationale	Selected for the customer and operational decision authority class, retail electricity context and the stated data, human-oversight, dependency and consequence profile.

Field	Content
Evidence	System risk assessment; architecture and authority record; representative validation; approval; monitoring; override, fallback and incident evidence.
Fallback	Revert to an approved deterministic procedure, trusted prior version or manual operation without bypassing independent protection.
Residual Risk	Accepted by the accountable owner after security, operations and applicable safety/reliability review.
Notably Absent	No claim that the use of AI improves reliability, safety, fairness or efficiency without system-specific operational evidence.

UC-15 - OT cyber anomaly detection

Field	Content
Subsector	Cross-sector
Authority Level	Security decision support
Objective	Use AI to support ot cyber anomaly detection while preserving accountable, bounded and recoverable operation.
Data Inputs	Telemetry, asset, operational, environmental, market, customer or security data as applicable; lineage and freshness are recorded.
Human Decision Point	Named operator or accountable owner reviews or authorises material action according to authority class.
Failure Modes	Incorrect prediction; stale, missing or manipulated data; distribution shift; dependency loss; automation bias; unauthorised or unsafe action.
Applicable Controls	D1-CTL-01; D1-CTL-03; D2-CTL-01; D2-CTL-02; D2-CTL-03; D2-CTL-04; D4-CTL-01; D4-CTL-02; D6-CTL-02; D6-CTL-04; D6-CTL-05; D8-CTL-03
Mapping Rationale	Selected for the security decision support authority class, cross-sector context and the stated data, human-oversight, dependency and consequence profile.
Evidence	System risk assessment; architecture and authority record; representative validation; approval; monitoring; override, fallback and incident evidence.
Fallback	Revert to an approved deterministic procedure, trusted prior version or manual operation without bypassing independent protection.
Residual Risk	Accepted by the accountable owner after security, operations and applicable safety/reliability review.
Notably Absent	No claim that the use of AI improves reliability, safety, fairness or efficiency without system-specific operational evidence.

UC-16 - Field-service decision support

Field	Content
Subsector	Cross-sector
Authority Level	Operational advisory
Objective	Use AI to support field-service decision support while preserving accountable, bounded and recoverable operation.
Data Inputs	Telemetry, asset, operational, environmental, market, customer or security data as applicable; lineage and freshness are recorded.
Human Decision Point	Named operator or accountable owner reviews or authorises material action according to authority class.
Failure Modes	Incorrect prediction; stale, missing or manipulated data; distribution shift; dependency loss; automation bias; unauthorised or unsafe action.
Applicable Controls	D1-CTL-01; D1-CTL-03; D4-CTL-02; D5-CTL-01; D6-CTL-01; D6-CTL-02; D7-CTL-01; D7-CTL-04; D8-CTL-01
Mapping Rationale	Selected for the operational advisory authority class, cross-sector context and the stated data, human-oversight,

Field	Content
	dependency and consequence profile.
Evidence	System risk assessment; architecture and authority record; representative validation; approval; monitoring; override, fallback and incident evidence.
Fallback	Revert to an approved deterministic procedure, trusted prior version or manual operation without bypassing independent protection.
Residual Risk	Accepted by the accountable owner after security, operations and applicable safety/reliability review.
Notably Absent	No claim that the use of AI improves reliability, safety, fairness or efficiency without system-specific operational evidence.

UC-17 - Billing anomaly detection

Field	Content
Subsector	Retail energy
Authority Level	Enterprise decision support
Objective	Use AI to support billing anomaly detection while preserving accountable, bounded and recoverable operation.
Data Inputs	Telemetry, asset, operational, environmental, market, customer or security data as applicable; lineage and freshness are recorded.
Human Decision Point	Named operator or accountable owner reviews or authorises material action according to authority class.
Failure Modes	Incorrect prediction; stale, missing or manipulated data; distribution shift; dependency loss; automation bias; unauthorised or unsafe action.
Applicable Controls	D1-CTL-01; D1-CTL-03; D5-CTL-01; D5-CTL-02; D5-CTL-05; D6-CTL-01; D6-CTL-02; D6-CTL-04; D8-CTL-01; D8-CTL-02
Mapping Rationale	Selected for the enterprise decision support authority class, retail energy context and the stated data, human-oversight, dependency and consequence profile.
Evidence	System risk assessment; architecture and authority record; representative validation; approval; monitoring; override, fallback and incident evidence.
Fallback	Revert to an approved deterministic procedure, trusted prior version or manual operation without bypassing independent protection.
Residual Risk	Accepted by the accountable owner after security, operations and applicable safety/reliability review.
Notably Absent	No claim that the use of AI improves reliability, safety, fairness or efficiency without system-specific operational evidence.

UC-18 - Physical security analytics

Field	Content
Subsector	Cross-sector
Authority Level	Security advisory
Objective	Use AI to support physical security analytics while preserving accountable, bounded and recoverable operation.
Data Inputs	Telemetry, asset, operational, environmental, market, customer or security data as applicable; lineage and freshness are recorded.
Human Decision Point	Named operator or accountable owner reviews or authorises material action according to authority class.
Failure Modes	Incorrect prediction; stale, missing or manipulated data; distribution shift; dependency loss; automation bias; unauthorised or unsafe action.
Applicable Controls	D1-CTL-01; D1-CTL-03; D2-CTL-05; D4-CTL-02; D5-CTL-01; D5-CTL-05; D6-CTL-01; D6-CTL-02; D7-CTL-04; D8-CTL-01; D9-CTL-03
Mapping Rationale	Selected for the security advisory authority class, cross-sector context and the stated data, human-oversight, dependency and

Field	Content
	consequence profile.
Evidence	System risk assessment; architecture and authority record; representative validation; approval; monitoring; override, fallback and incident evidence.
Fallback	Revert to an approved deterministic procedure, trusted prior version or manual operation without bypassing independent protection.
Residual Risk	Accepted by the accountable owner after security, operations and applicable safety/reliability review.
Notably Absent	No claim that the use of AI improves reliability, safety, fairness or efficiency without system-specific operational evidence.

Composite worked scenarios

CS-01 - Real-time DER dispatch for voltage support

Field	Content
Systems involved	Load and renewable forecasts; battery optimisation; DER orchestration; DMS/SCADA and field controllers.
Authority and control flow	Forecasts produce a proposed dispatch; operational constraints and safety checks validate it; material actions require approved authority; commands are signed and reconciled; independent protection remains active.
Connected controls	D1-CTL-01; D1-CTL-03; D3-CTL-01; D3-CTL-02; D4-CTL-02; D6-CTL-01; D6-CTL-02; D6-CTL-04; D6-CTL-05; D9-CTL-01; D9-CTL-02; D9-CTL-03; D9-CTL-06
Failure and stress tests	Manipulated forecast, aggregator loss, stale telemetry, correlated DER response, unsafe rate of change and failed rollback.
Fallback	Local voltage controls and operator-approved deterministic dispatch.

CS-02 - AI-assisted outage restoration

Field	Content
Systems involved	Damage assessment; outage prediction; critical-load records; crew/asset availability; restoration planning.
Authority and control flow	AI proposes priorities; field and critical-service data are reconciled; an authorised restoration lead approves; execution and deviations are logged.
Connected controls	D1-CTL-01; D1-CTL-03; D5-CTL-01; D5-CTL-02; D6-CTL-01; D6-CTL-02; D6-CTL-04; D7-CTL-04; D8-CTL-01; D9-CTL-02
Failure and stress tests	Incomplete field data, biased prioritisation, communications loss, critical dependency omission and operator overload.
Fallback	Approved manual emergency-restoration plan.

CS-03 - Pipeline anomaly detection and response

Field	Content
Systems involved	Sensor and historian data; hydraulic model; anomaly detection; control-room workflow; field inspection.
Authority and control flow	AI raises an anomaly with confidence and data-quality status; operators correlate independent evidence; shutdown or isolation remains governed by approved safety and operating procedures.
Connected controls	D1-CTL-01; D1-CTL-03; D1-CTL-08; D2-CTL-04; D4-CTL-02; D6-CTL-02; D6-CTL-04; D9-CTL-01; D9-CTL-02; D9-CTL-03
Failure and stress tests	Suppressed alert, false sensor data, stale telemetry, threshold tampering and cloud/communications loss.
Fallback	Independent leak detection, field verification and established

Field	Content
	emergency procedure.

CS-04 - Automated energy-market bidding

Field	Content
Systems involved	Forecasts; portfolio and risk data; trading agent; market gateway; compliance surveillance.
Authority and control flow	The agent generates bounded orders; policy and market-rule checks apply; threshold breaches require approval; orders and rationale are recorded and monitored.
Connected controls	D1-CTL-01; D1-CTL-03; D2-CTL-01; D3-CTL-01; D3-CTL-02; D5-CTL-01; D5-CTL-02; D5-CTL-04; D6-CTL-01; D6-CTL-02; D6-CTL-04; D7-CTL-05; D8-CTL-01; D8-CTL-02; D8-CTL-04
Failure and stress tests	Objective manipulation, abnormal bids, compromised credentials, stale market data and prohibited coordination.
Fallback	Suspend automation and use approved manual trading controls.

Threat, failure and misuse scenarios

SC-01 - Manipulated load or generation forecasts

Field	Content
Category	Adversarial, accidental or operational failure; classify from evidence rather than assumption.
Initiating Condition	Adversarial data manipulation or compromised upstream data.
Affected Assets	Models, data, cloud/edge platforms, OT interfaces, operators, field devices, safety/reliability functions and business processes as applicable.
Consequence	Potential safety, reliability, availability, market, financial, customer, environmental or regulatory harm, including cascading effects.
Preventive Controls	Authenticate and validate data sources; diversify critical feeds; monitor provenance and forecast sensitivity.
Detective Controls	Compare independent forecasts; detect residual, topology and source anomalies; alert on abrupt confidence changes.
Recovery Controls	Isolate affected feeds; revert to trusted or conservative forecast; require operator review; preserve data and model evidence.
Response Owner	Named incident commander with operations, OT security, engineering, safety/reliability, legal/compliance and supplier support as applicable.
Residual Risk	Record after controls, exercises and accepted limitations.
Confidence	Scenario analysis; frequency and magnitude require local evidence.
Notably Absent	The scenario does not establish that any observed event was caused by AI or by a hostile actor.

SC-02 - Compromised model-update pipeline

Field	Content
Category	Adversarial, accidental or operational failure; classify from evidence rather than assumption.
Initiating Condition	Supply-chain compromise, malicious artefact substitution or unauthorised release.
Affected Assets	Models, data, cloud/edge platforms, OT interfaces, operators, field devices, safety/reliability functions and business processes as applicable.
Consequence	Potential safety, reliability, availability, market, financial, customer, environmental or regulatory harm, including cascading effects.

Field	Content
Preventive Controls	Signed artefacts; segregated build/release roles; provenance; reproducible build; approved update and rollback process.
Detective Controls	Signature/hash mismatch; unexpected dependency/model change; behavioural regression; unauthorised release event.
Recovery Controls	Stop deployment; revoke credentials; restore trusted version; investigate build chain; notify affected parties and suppliers.
Response Owner	Named incident commander with operations, OT security, engineering, safety/reliability, legal/compliance and supplier support as applicable.
Residual Risk	Record after controls, exercises and accepted limitations.
Confidence	Scenario analysis; frequency and magnitude require local evidence.
Notably Absent	The scenario does not establish that any observed event was caused by AI or by a hostile actor.

SC-03 - False sensor data influencing AI recommendations

Field	Content
Category	Adversarial, accidental or operational failure; classify from evidence rather than assumption.
Initiating Condition	Sensor compromise, calibration failure, replay or telemetry corruption.
Affected Assets	Models, data, cloud/edge platforms, OT interfaces, operators, field devices, safety/reliability functions and business processes as applicable.
Consequence	Potential safety, reliability, availability, market, financial, customer, environmental or regulatory harm, including cascading effects.
Preventive Controls	Sensor authentication; plausibility and redundancy checks; time synchronisation; secure gateways and data-quality rules.
Detective Controls	Cross-sensor inconsistency; impossible rate-of-change; stale/replayed timestamps; divergence from physical state.
Recovery Controls	Quarantine source; switch to validated redundant/manual inputs; constrain AI authority; inspect field device and preserve telemetry.
Response Owner	Named incident commander with operations, OT security, engineering, safety/reliability, legal/compliance and supplier support as applicable.
Residual Risk	Record after controls, exercises and accepted limitations.
Confidence	Scenario analysis; frequency and magnitude require local evidence.
Notably Absent	The scenario does not establish that any observed event was caused by AI or by a hostile actor.

SC-04 - Topology-model corruption

Field	Content
Category	Adversarial, accidental or operational failure; classify from evidence rather than assumption.
Initiating Condition	Incorrect, stale or malicious network/process topology data.
Affected Assets	Models, data, cloud/edge platforms, OT interfaces, operators, field devices, safety/reliability functions and business processes as applicable.
Consequence	Potential safety, reliability, availability, market, financial, customer, environmental or regulatory harm, including cascading effects.
Preventive Controls	Controlled topology changes; dual approval; signed versions; reconciliation with operational configuration.
Detective Controls	Mismatch between model and observed switching/flow state; failed state estimation; unexpected connectivity.
Recovery Controls	Freeze automated optimisation; restore verified topology; reconcile field state; rerun safety and contingency checks.

Field	Content
Response Owner	Named incident commander with operations, OT security, engineering, safety/reliability, legal/compliance and supplier support as applicable.
Residual Risk	Record after controls, exercises and accepted limitations.
Confidence	Scenario analysis; frequency and magnitude require local evidence.
Notably Absent	The scenario does not establish that any observed event was caused by AI or by a hostile actor.

SC-05 - AI action outside authorised operating bounds

Field	Content
Category	Adversarial, accidental or operational failure; classify from evidence rather than assumption.
Initiating Condition	Defective policy, compromise, configuration error or emergent agent behaviour.
Affected Assets	Models, data, cloud/edge platforms, OT interfaces, operators, field devices, safety/reliability functions and business processes as applicable.
Consequence	Potential safety, reliability, availability, market, financial, customer, environmental or regulatory harm, including cascading effects.
Preventive Controls	Action allow-list; operating envelope; rate limits; dual approval; independent interlocks; least agency.
Detective Controls	Pre-execution policy rejection; interlock activation; post-command deviation; abnormal action frequency.
Recovery Controls	Block or abort action; enter safe state; transfer to manual/deterministic control; revoke machine authority; investigate.
Response Owner	Named incident commander with operations, OT security, engineering, safety/reliability, legal/compliance and supplier support as applicable.
Residual Risk	Record after controls, exercises and accepted limitations.
Confidence	Scenario analysis; frequency and magnitude require local evidence.
Notably Absent	The scenario does not establish that any observed event was caused by AI or by a hostile actor.

SC-06 - Operator over-reliance on incorrect recommendations

Field	Content
Category	Adversarial, accidental or operational failure; classify from evidence rather than assumption.
Initiating Condition	Automation bias, poor uncertainty display, workload or inadequate training.
Affected Assets	Models, data, cloud/edge platforms, OT interfaces, operators, field devices, safety/reliability functions and business processes as applicable.
Consequence	Potential safety, reliability, availability, market, financial, customer, environmental or regulatory harm, including cascading effects.
Preventive Controls	Human-factors design; confidence/limitations display; training; challenge and dissent procedures; workload assessment.
Detective Controls	Low override rate despite poor performance; repeated acceptance without review; near misses; user feedback.
Recovery Controls	Suspend recommendation use; require secondary verification; retrain operators; redesign interface and approval procedure.
Response Owner	Named incident commander with operations, OT security, engineering, safety/reliability, legal/compliance and supplier support as applicable.
Residual Risk	Record after controls, exercises and accepted limitations.
Confidence	Scenario analysis; frequency and magnitude require local

Field	Content
	evidence.
Notably Absent	The scenario does not establish that any observed event was caused by AI or by a hostile actor.

SC-07 - Loss of cloud connectivity

Field	Content
Category	Adversarial, accidental or operational failure; classify from evidence rather than assumption.
Initiating Condition	Telecommunications failure, provider outage, route disruption or account suspension.
Affected Assets	Models, data, cloud/edge platforms, OT interfaces, operators, field devices, safety/reliability functions and business processes as applicable.
Consequence	Potential safety, reliability, availability, market, financial, customer, environmental or regulatory harm, including cascading effects.
Preventive Controls	Local fallback; offline data/cache policy; dependency inventory; tested failover; no sole dependence for safety function.
Detective Controls	Heartbeat and latency monitoring; stale-data alarms; failed API calls; dependency health alerts.
Recovery Controls	Fail locally to deterministic/manual control; prevent stale recommendations; preserve local logs; recover and reconcile later.
Response Owner	Named incident commander with operations, OT security, engineering, safety/reliability, legal/compliance and supplier support as applicable.
Residual Risk	Record after controls, exercises and accepted limitations.
Confidence	Scenario analysis; frequency and magnitude require local evidence.
Notably Absent	The scenario does not establish that any observed event was caused by AI or by a hostile actor.

SC-08 - Model drift after grid or plant reconfiguration

Field	Content
Category	Adversarial, accidental or operational failure; classify from evidence rather than assumption.
Initiating Condition	Topology, equipment, process, season or operating-regime change.
Affected Assets	Models, data, cloud/edge platforms, OT interfaces, operators, field devices, safety/reliability functions and business processes as applicable.
Consequence	Potential safety, reliability, availability, market, financial, customer, environmental or regulatory harm, including cascading effects.
Preventive Controls	Change-triggered reassessment; shadow mode; representative retraining/validation; configuration linkage.
Detective Controls	Performance degradation by regime; residual shift; OOD alerts; increasing overrides or constraint violations.
Recovery Controls	Reduce or suspend authority; use trusted baseline/manual method; revalidate with new configuration before restoration.
Response Owner	Named incident commander with operations, OT security, engineering, safety/reliability, legal/compliance and supplier support as applicable.
Residual Risk	Record after controls, exercises and accepted limitations.
Confidence	Scenario analysis; frequency and magnitude require local evidence.
Notably Absent	The scenario does not establish that any observed event was caused by AI or by a hostile actor.

SC-09 - Compromised vendor remote support

Field	Content
Category	Adversarial, accidental or operational failure; classify from evidence rather than assumption.
Initiating Condition	Stolen credentials, supplier compromise or unauthorised remote activity.
Affected Assets	Models, data, cloud/edge platforms, OT interfaces, operators, field devices, safety/reliability functions and business processes as applicable.
Consequence	Potential safety, reliability, availability, market, financial, customer, environmental or regulatory harm, including cascading effects.
Preventive Controls	Just-in-time access; MFA; session approval/recording; segmentation; supplier identity and incident clauses.
Detective Controls	Unexpected access time/location; privileged commands; configuration changes; session anomalies.
Recovery Controls	Terminate session; revoke credentials; isolate pathway; restore configurations; investigate supplier and notify.
Response Owner	Named incident commander with operations, OT security, engineering, safety/reliability, legal/compliance and supplier support as applicable.
Residual Risk	Record after controls, exercises and accepted limitations.
Confidence	Scenario analysis; frequency and magnitude require local evidence.
Notably Absent	The scenario does not establish that any observed event was caused by AI or by a hostile actor.

SC-10 - Coordinated manipulation of distributed energy resources

Field	Content
Category	Adversarial, accidental or operational failure; classify from evidence rather than assumption.
Initiating Condition	Compromised aggregator, device fleet, command channel or market signal.
Affected Assets	Models, data, cloud/edge platforms, OT interfaces, operators, field devices, safety/reliability functions and business processes as applicable.
Consequence	Potential safety, reliability, availability, market, financial, customer, environmental or regulatory harm, including cascading effects.
Preventive Controls	Fleet segmentation; command signing; per-device/fleet limits; diversity; aggregator assurance; anti-replay.
Detective Controls	Correlated abnormal response; command/telemetry mismatch; geographic or device-cluster anomalies.
Recovery Controls	Block aggregator commands; constrain fleet response; island affected groups; coordinate operator/market response; preserve evidence.
Response Owner	Named incident commander with operations, OT security, engineering, safety/reliability, legal/compliance and supplier support as applicable.
Residual Risk	Record after controls, exercises and accepted limitations.
Confidence	Scenario analysis; frequency and magnitude require local evidence.
Notably Absent	The scenario does not establish that any observed event was caused by AI or by a hostile actor.

SC-11 - Unsafe battery optimisation

Field	Content
Category	Adversarial, accidental or operational failure; classify from evidence rather than assumption.
Initiating Condition	Incorrect state estimation, objective function, constraints or

Field	Content
	malicious command.
Affected Assets	Models, data, cloud/edge platforms, OT interfaces, operators, field devices, safety/reliability functions and business processes as applicable.
Consequence	Potential safety, reliability, availability, market, financial, customer, environmental or regulatory harm, including cascading effects.
Preventive Controls	Independent BMS protections; SOC/SOH validation; thermal and power envelopes; command rate limits; safe-state logic.
Detective Controls	Protection activation; thermal anomaly; SOC inconsistency; cycling outside approved limits; command rejection.
Recovery Controls	Stop optimisation; defer to BMS/local controller; isolate asset if required; inspect model, sensors and constraints.
Response Owner	Named incident commander with operations, OT security, engineering, safety/reliability, legal/compliance and supplier support as applicable.
Residual Risk	Record after controls, exercises and accepted limitations.
Confidence	Scenario analysis; frequency and magnitude require local evidence.
Notably Absent	The scenario does not establish that any observed event was caused by AI or by a hostile actor.

SC-12 - Pipeline anomaly-detection suppression

Field	Content
Category	Adversarial, accidental or operational failure; classify from evidence rather than assumption.
Initiating Condition	Poisoning, alert threshold manipulation, sensor compromise or analyst override misuse.
Affected Assets	Models, data, cloud/edge platforms, OT interfaces, operators, field devices, safety/reliability functions and business processes as applicable.
Consequence	Potential safety, reliability, availability, market, financial, customer, environmental or regulatory harm, including cascading effects.
Preventive Controls	Protected thresholds; independent leak/protection layers; provenance; role separation; secure configuration.
Detective Controls	Alert-rate discontinuity; unexplained threshold changes; discrepancy with hydraulic/field observations.
Recovery Controls	Reinstate validated thresholds; use independent detection and field inspection; isolate model; preserve and report evidence.
Response Owner	Named incident commander with operations, OT security, engineering, safety/reliability, legal/compliance and supplier support as applicable.
Residual Risk	Record after controls, exercises and accepted limitations.
Confidence	Scenario analysis; frequency and magnitude require local evidence.
Notably Absent	The scenario does not establish that any observed event was caused by AI or by a hostile actor.

SC-13 - Automated bidding market manipulation

Field	Content
Category	Adversarial, accidental or operational failure; classify from evidence rather than assumption.
Initiating Condition	Objective misalignment, compromised agent, collusion signal or rule violation.
Affected Assets	Models, data, cloud/edge platforms, OT interfaces, operators, field devices, safety/reliability functions and business processes as applicable.
Consequence	Potential safety, reliability, availability, market, financial, customer, environmental or regulatory harm, including

Field	Content
	cascading effects.
Preventive Controls	Bounded strategies; market-rule constraints; approval thresholds; conflict controls; immutable decision logs.
Detective Controls	Abnormal bid patterns; unexplained strategy shifts; concentration or coordination indicators; rejected orders.
Recovery Controls	Suspend automated bidding; cancel/limit orders where permitted; switch to approved manual process; legal/compliance review.
Response Owner	Named incident commander with operations, OT security, engineering, safety/reliability, legal/compliance and supplier support as applicable.
Residual Risk	Record after controls, exercises and accepted limitations.
Confidence	Scenario analysis; frequency and magnitude require local evidence.
Notably Absent	The scenario does not establish that any observed event was caused by AI or by a hostile actor.

SC-14 - AI cyber-response disrupts legitimate operations

Field	Content
Category	Adversarial, accidental or operational failure; classify from evidence rather than assumption.
Initiating Condition	False positive, excessive containment authority or adversarial trigger.
Affected Assets	Models, data, cloud/edge platforms, OT interfaces, operators, field devices, safety/reliability functions and business processes as applicable.
Consequence	Potential safety, reliability, availability, market, financial, customer, environmental or regulatory harm, including cascading effects.
Preventive Controls	Least agency; approved response playbooks; asset criticality context; simulation; human approval for disruptive actions.
Detective Controls	High false-positive rate; containment of critical assets; operator conflict; unusual response sequence.
Recovery Controls	Stop autonomous response; restore connectivity/configuration safely; use manual incident command; review detection and authority.
Response Owner	Named incident commander with operations, OT security, engineering, safety/reliability, legal/compliance and supplier support as applicable.
Residual Risk	Record after controls, exercises and accepted limitations.
Confidence	Scenario analysis; frequency and magnitude require local evidence.
Notably Absent	The scenario does not establish that any observed event was caused by AI or by a hostile actor.

SC-15 - Restoration prioritisation failure

Field	Content
Category	Adversarial, accidental or operational failure; classify from evidence rather than assumption.
Initiating Condition	Bad damage data, biased objective, communications loss or incorrect dependency model.
Affected Assets	Models, data, cloud/edge platforms, OT interfaces, operators, field devices, safety/reliability functions and business processes as applicable.
Consequence	Potential safety, reliability, availability, market, financial, customer, environmental or regulatory harm, including cascading effects.
Preventive Controls	Critical-load and dependency validation; human approval; alternative plans; scenario exercises; equity/service constraints.
Detective Controls	Plan infeasibility; conflict with field reports; repeated

Field	Content
	reprioritisation; omission of critical dependencies.
Recovery Controls	Switch to operator-led restoration procedure; validate critical services; reconcile field status; document decisions.
Response Owner	Named incident commander with operations, OT security, engineering, safety/reliability, legal/compliance and supplier support as applicable.
Residual Risk	Record after controls, exercises and accepted limitations.
Confidence	Scenario analysis; frequency and magnitude require local evidence.
Notably Absent	The scenario does not establish that any observed event was caused by AI or by a hostile actor.

SC-16 - Shared model defect across operators

Field	Content
Category	Adversarial, accidental or operational failure; classify from evidence rather than assumption.
Initiating Condition	Common supplier defect, common dataset error or correlated update.
Affected Assets	Models, data, cloud/edge platforms, OT interfaces, operators, field devices, safety/reliability functions and business processes as applicable.
Consequence	Potential safety, reliability, availability, market, financial, customer, environmental or regulatory harm, including cascading effects.
Preventive Controls	Diversity for critical functions; staged rollout; supplier concentration analysis; coordinated disclosure terms.
Detective Controls	Correlated anomalies across sites/operators; common version fingerprint; sector information-sharing alerts.
Recovery Controls	Stop rollout; revert affected estates; coordinate with supplier and sector partners; assess systemic exposure.
Response Owner	Named incident commander with operations, OT security, engineering, safety/reliability, legal/compliance and supplier support as applicable.
Residual Risk	Record after controls, exercises and accepted limitations.
Confidence	Scenario analysis; frequency and magnitude require local evidence.
Notably Absent	The scenario does not establish that any observed event was caused by AI or by a hostile actor.

SC-17 - Rollback failure after model release

Field	Content
Category	Adversarial, accidental or operational failure; classify from evidence rather than assumption.
Initiating Condition	Incompatible data/schema, missing prior artefact, untested rollback or state corruption.
Affected Assets	Models, data, cloud/edge platforms, OT interfaces, operators, field devices, safety/reliability functions and business processes as applicable.
Consequence	Potential safety, reliability, availability, market, financial, customer, environmental or regulatory harm, including cascading effects.
Preventive Controls	Versioned artefacts and schemas; rollback rehearsal; state migration plan; compatibility and backup tests.
Detective Controls	Rollback test failure; checksum mismatch; incompatible state; inability to restore prior performance.
Recovery Controls	Enter deterministic/manual mode; restore backup and compatible dependencies; rebuild trusted environment; investigate release gate.
Response Owner	Named incident commander with operations, OT security, engineering, safety/reliability, legal/compliance and supplier support as applicable.

Field	Content
	support as applicable.
Residual Risk	Record after controls, exercises and accepted limitations.
Confidence	Scenario analysis; frequency and magnitude require local evidence.
Notably Absent	The scenario does not establish that any observed event was caused by AI or by a hostile actor.

SC-18 - Training-data leakage of critical infrastructure details

Field	Content
Category	Adversarial, accidental or operational failure; classify from evidence rather than assumption.
Initiating Condition	Unauthorised training use, memorisation, extraction, exposure or supplier misuse.
Affected Assets	Models, data, cloud/edge platforms, OT interfaces, operators, field devices, safety/reliability functions and business processes as applicable.
Consequence	Potential safety, reliability, availability, market, financial, customer, environmental or regulatory harm, including cascading effects.
Preventive Controls	Data classification/minimisation; approved training purpose; access control; contractual restrictions; privacy/security testing.
Detective Controls	Extraction test; unusual model disclosure; access audit; supplier notification; sensitive-content scan.
Recovery Controls	Restrict access and revoke credentials; contain model/version; assess deletion or retraining/unlearning feasibility; legal, security and affected-party notification as applicable.
Response Owner	Named incident commander with operations, OT security, engineering, safety/reliability, legal/compliance and supplier support as applicable.
Residual Risk	Record after controls, exercises and accepted limitations.
Confidence	Scenario analysis; frequency and magnitude require local evidence.
Notably Absent	The scenario does not establish that any observed event was caused by AI or by a hostile actor.

Assessment and evidence model

Source Verification and Assurance Maturity are independent.

Scale	Definition
T1	Direct primary or authoritative evidence.
T2	Strong corroborated evidence.
T3	Limited or indirect evidence.
T4	Unverified, disputed or insufficient evidence.
A0	Not established.
A1	Ad hoc.
A2	Defined.
A3	Implemented.
A4	Measured.
A5	Independently assured and continuously improved.

Notably Absent

- No universal evidence that AI improves energy reliability, safety, fairness, efficiency or operator workload.
- No proof that accuracy, simulation, a human-in-the-loop label, a supplier attestation or a provenance mechanism establishes safe operation.

- No automatic equivalence between GAISSF implementation and legal, regulatory, safety, nuclear or certification compliance.
- No universal architecture, threshold, maturity level or F/O/A sequencing label applies to every energy system.
- No assumption that every anomaly, outage or unsafe outcome is AI-related or hostile.
- No claim that the differentiated use-case mappings or scenario treatments have completed specialist validation.

Source register

ID	Source	Issuer	Status	Use	Limitation
SRC-001	GAISSF-NOR-001 Framework Standard v1.0	ODA3 Institute	Corrected Final Publication Edition	Control identifiers, titles, domains and conformance language.	Sector implementation remains informative.
SRC-002	Artificial Intelligence Risk Management Framework (AI RMF 1.0), NIST AI 100-1	NIST	Published	AI risk-management lifecycle context.	Does not establish energy-sector compliance.
SRC-003	Regulation (EU) 2024/1689 (Artificial Intelligence Act)	European Union	In force; phased application	EU AI legal context where applicable.	Applicability and timelines require qualified legal review.
SRC-004	Cybersecurity - Office of Cybersecurity, Energy Security, and Emergency Response	U.S. Department of Energy	Current portal; individual materials vary	Energy cybersecurity context.	Portal content is not a single compliance standard.
SRC-005	Guide to Operational Technology Security, NIST SP 800-82 Rev. 3	NIST	Final	OT security architecture and operational constraints.	Not AI-specific.
SRC-006	IEC 62443 series - Security for industrial automation and control systems	IEC	Published series	Industrial automation and control-system security context.	Licensed text; part, edition and applicability must be verified.
SRC-007	ISO/IEC 42001 - Artificial intelligence management system	ISO/IEC	Published	AI management-system context.	Mapping does not establish equivalence or certification.
SRC-008	ISO/IEC 23894 - Artificial intelligence - Guidance on risk management	ISO/IEC	Published	AI risk-management context.	Edition and corrigenda must be verified.
SRC-009	Regulation (EU) 2022/2554 on digital operational resilience for the financial sector (DORA)	European Union	Applicable from 2025-01-17	Context for the authoritative D8-CTL-04 title and financial-sector operational-resilience obligations where legally in scope.	DORA is not an energy-sector law merely because the GAISSF control title references it; apply only to legally covered entities.
SRC-010	NIST SP 800-218A - Secure Software Development Practices for Generative AI and Dual-Use Foundation Models: An SSDF Community Profile	NIST	Final	AI-specific secure-development practices.	Use with NIST SP 800-218; applicability to non-generative systems must be assessed.
SRC-011	IEC 61508 series - Functional safety of electrical/electronic/programmable electronic safety-related systems	IEC	Published series	Functional-safety lifecycle and independent protection context.	Licensed text; sector-specific derivatives and competent safety interpretation may take precedence.
SRC-012	DO-178C - Software Considerations in Airborne Systems and Equipment Certification	RTCA	Published	Illustrative high-assurance software evidence context only where aviation-related energy systems are in scope.	Not an energy-sector standard and not a basis for general energy compliance.
SRC-013	Open Security	NIST	Published project	Optional machine-	Use does not prove

ID	Source	Issuer	Status	Use	Limitation
	Controls Assessment Language (OSCAL)			readable control, assessment and evidence exchange.	control effectiveness or conformance.
SRC-014	C2PA Technical Specification	Coalition for Content Provenance and Authenticity	Published specification	Optional content-provenance context.	Provenance metadata does not establish truth, safety or fitness for high-consequence operation.

Publication gates

Gate	Review	Status	Closure evidence
PG-01	Control reconciliation	Passed	Automated and manual count/title/domain reconciliation against GAISSF-NOR-001.
PG-02	Energy-sector technical review	Open	Qualified energy-operations review.
PG-03	OT/ICS security review	Open	Qualified OT/ICS security review.
PG-04	Power-system or process-engineering review	Open	Qualified engineering review appropriate to covered subsectors.
PG-05	Safety review	Open	Qualified functional/process/electrical safety review.
PG-06	AI security review	Open	Independent AI security review.
PG-07	Regulatory mapping review	Open	Jurisdiction-specific regulatory review.
PG-08	Legal review	Open	Qualified legal review.
PG-09	Privacy review	Open	Privacy and data-protection review where applicable.
PG-10	Accessibility review	Open	Accessibility audit and remediation record.
PG-11	Editorial review	Open	Final editorial and house-style approval.
PG-12	Workbook QA	Passed with Conditions	Structural, formula-error and render checks passed; adopting organisations must populate templates.
PG-13	CSV validation	Passed	RFC 4180 round-trip, 59 rows, stable columns and formula-injection scan.
PG-14	JSON validation	Passed	Strict parse, schema-shape and recursive whitespace checks.
PG-15	PDF pagination review	Passed	All pages rendered and visually inspected for clipping and overlap.
PG-16	Hyperlink review	Passed with Conditions	Visible source locators checked; external availability may change.
PG-17	Checksum and manifest reconciliation	Passed	Payload hashes and manifest records reconciled.
PG-18	Licence review	Open	Approved publication licence and trademark wording review.
PG-19	Publication approval	Open	Formal authorised publication approval.
PG-20	Use-case control-mapping validation	Open	Energy, OT, safety and GAISSF specialists validate differentiated mappings.
PG-21	Scenario stress-test review	Open	Scenario-specific preventive, detective and recovery treatments exercised or reviewed.

Validation statement

Structural validation confirms counts, identifier uniqueness, parseability, schema shape, CSV round-trip and cross-file dataset parity. Semantic correctness of use-case mappings and scenario treatments remains subject to PG-20 and PG-21 specialist review; structural validation is not represented as semantic approval.