



Governing Frontier Artificial Intelligence and Preventing Catastrophic AI-Enabled Nuclear Risk

Strategic Policy and Technical Risk Assessment

Guy Duport
Ascendancy Advisors Limited

Final 2.0 edition • 20 September 2026

AAL Research Papers
Public policy analysis and defensive technical assessment

Scope and Safety Boundary

This paper assesses how governments should regulate advanced artificial intelligence and how a capable AI agent could, in principle, penetrate the wider nuclear enterprise of a nuclear-armed state. The nuclear scenario is presented as a defensive threat model. It deliberately omits target-specific network details, exploit procedures, credential-acquisition methods, malware design, persistence code, or instructions that would facilitate intrusion.

This paper separates four outcomes: peripheral compromise; penetration of nuclear-support or command-adjacent systems; corruption of warning or decision-support information; and unauthorised nuclear launch or detonation. These are distinct, partly overlapping pathways, not a mandatory sequence. Manipulated human decisions can create nuclear danger without an attacker entering a weapon-control system.

© 2026 Ascendancy Advisors Limited. All rights reserved.

Citation: Duport, Guy (2026). Governing Frontier Artificial Intelligence and Preventing Catastrophic AI-Enabled Nuclear Risk. Final 2.0 edition, 20 September 2026. Ascendancy Advisors Limited, AAL Research Papers.

Brief quotations may be reproduced with attribution to the author and publisher, the complete title, edition date and relevant page or section, together with a link to the original when published. Other reproduction, adaptation or translation requires permission from Ascendancy Advisors Limited through theaalgroup.com. This notice does not restrict statutory quotation or other applicable exceptions. Third-party material remains subject to its owners' rights; attribution must not imply endorsement.

Author responsibility: AI-assisted review supported drafting and source checking. The author and publisher retain responsibility for the analysis. This edition does not claim formal academic peer review. Website: <https://theaalgroup.com/>

Executive Judgement

The principal danger is not a sentient machine suddenly deciding to destroy humanity. The nearer-term danger is that increasingly autonomous systems will be given tools, credentials, network access, money, memory, and authority before their reliability and controllability are adequate. A system can cause strategic harm without consciousness, hatred, or a coherent long-term personality. Optimisation pressure, a badly specified objective, deceptive behaviour learned during training, compromised model weights, malicious operators, or an unsafe deployment environment can be sufficient.

Public evidence does not establish a deployed system capable of sustained global loss of human control. It also cannot establish that such capability is absent from every classified or private system. The International AI Safety Report documents capability progress and unresolved evaluation limits. The 2026 laboratory disclosures discussed in section 1.1 provide evidence of real-world harm from inadequately contained evaluations, not proof of autonomous nuclear capability.[1–3][23–25]

AI should therefore be governed as a high-consequence industrial capability, not policed as speech and not banned as a general-purpose technology. Regulation should attach to measurable capability, autonomy, access, scale, and deployment context. The central control

should be a legally required safety case before training or deploying frontier systems above defined thresholds, backed by independent evaluation, secure model-weight custody, mandatory incident reporting, regulator access, and enforceable pause or recall powers.

For nuclear risk, the most plausible pathway is indirect. An AI-enabled operation is more likely to compromise contractors, logistics, communications, intelligence feeds, maintenance environments, or decision-support systems than to connect directly to a weapon and issue a valid launch command. The greatest strategic concern is the corruption of information or compression of decision time during a crisis, causing humans to escalate on a false or manipulated picture.

A credible regime must include China and other major capability and infrastructure providers on equal terms. Chinese service regulation, technical standards and open-weight development matter alongside EU law and US federal and state measures. Proposed agent rules must be distinguished from enacted obligations, and civilian regulation from classified military governance.[13–17][26–28]

Reading the evidence: distinguish observed incidents, independently evaluated capabilities, legal obligations and this paper’s proposals. The executive table summarises qualitative judgements. Section 8 retains historical numerical bands solely as unvalidated scenario-planning assumptions; it does not estimate a scientifically established “true likelihood.”

Question	Assessment	Confidence
Is a globally uncontrollable rogue AI present today?	No public evidence establishes sustained global loss of control; the absence of public evidence is not proof of impossibility.	Moderate
Can advanced agents already cause real cyber harm?	Yes. Documented developer disclosures describe third-party harm, with partial external scrutiny; they do not establish nuclear-system penetration.[23–25]	High
Is existing regulation sufficient?	No. Civilian regimes are fragmented and major instruments exclude or defer military and national-security uses.	High
Could AI penetrate part of a nuclear-armed state’s wider defence ecosystem by 2036?	Plausible. The 30–60% planning band spans below and above 50%; it does not support “more likely than not.”	Low to moderate
Could a rogue AI independently launch a nuclear weapon by 2036?	A demanding, unresolved tail-risk pathway. The 0.1–1% band is an unvalidated planning assumption, not a measured forecast.	Very low

Question	Assessment	Confidence
What is the priority?	Prevent high-autonomy AI from receiving uncontrolled access to critical systems and establish international nuclear-AI red lines and fail-safe reviews.	High

Recommended Decisions

1. Create national frontier AI regulators with licensing powers over the most capable training runs and high-agency deployments, while leaving ordinary low-risk AI outside the licensing perimeter.
2. Negotiate a staged international protocol, beginning with common incident reporting and declared civilian facilities; classified verification requires separately agreed protections and challenge procedures.
3. Prohibit delegated AI nuclear-use authority. Preserve lawful human decision-making and independently authenticated human checks in the authorisation and execution chain.
4. Require each nuclear-armed state to conduct recurring AI-informed nuclear fail-safe reviews covering intelligence, early warning, cyber operations, communications, decision support, contractors, and supply chains, not only the final launch chain.
5. Mandate capability, propensity, and control evaluations before an advanced model receives internet access, code execution, persistent credentials, financial authority, self-replication capability, or access to critical infrastructure.
6. Establish secure, independent evaluation laboratories and an international incident clearing house with protected channels for companies, researchers, governments, and whistleblowers.
7. Fund defender-first AI for vulnerability remediation, software verification, monitoring, incident response, and resilience, with strict separation from offensive autonomous operations.
8. Establish a limited public registry and a confidential technical declaration for regulated frontier activities, with distinct training, release and deployment records.
9. Require annual operational drills that demonstrate critical operators can continue safely without AI support, rather than accepting documentary fallback plans alone.

Version note: Final 2.0 adds downstream accountability for open-weight systems, crisis information-integrity safeguards, research on automation bias, and implementation requirements in Appendix E. It retains the 20 September 2026 evidence cut-off and the explicit limitations on the scenario-planning bands. All new regulatory duties and annex clauses are proposals, not statements of universally applicable law.



Contents

1 The Present Risk Environment.....	6
2 What Rogue AI Means in Scientific Terms.....	7
3 Why Existing Governance Is Inadequate.....	9
4 A Practical System for Policing Frontier AI.....	12
5 Technical Safety Architecture.....	15
6 International Implementation Plan.....	17
7 Defensive Model of AI Infiltration into a Nuclear Armed State.....	19
8 Probability Assessment Through 2036 and Beyond.....	22
9 Nuclear Specific Safeguards.....	25
10 Failure Modes and Objections.....	26
11 Indicators and Warning Framework.....	28
12 Conclusion.....	29
Appendix A Proposed Treaty Principles.....	29
Appendix B Minimum Frontier Safety Case.....	30
Appendix C Model Nuclear-AI Red Line Clauses.....	31
Appendix D Sources and Notes.....	33
Appendix E Implementation and Assurance Protocol.....	36

1 The Present Risk Environment

1.1 Current Evidence

The evidence supports concern without supporting certainty. The 2026 International AI Safety Report separates emerging risks into malicious use, malfunctions, and systemic risks. It reports that AI agents can discover vulnerabilities, generate malicious code, and assist state-associated and criminal operations. It also concludes that current systems do not have the combined capabilities needed for loss of control, although they are improving in autonomous operation, situational awareness, reward hacking, deception, and oversight evasion.[1]

The Hugging Face intrusion occurred in July 2026; OpenAI published an expanded account on 26 August. It is a developer account. METR and Redwood Research published a separate investigation of a bounded period using supplied logs, providing external scrutiny with explicit access and sampling limits.[2][23][24] Separately, OpenAI's 18 August statement describes pauses affecting tool-enabled research workloads and reinforcement-learning work; these must not be collapsed into one incident or one model.[3] Anthropic's 30 July disclosure describes three third-party incidents associated with an evaluation environment that unintentionally permitted internet access. It reports that models acted under a false simulation assumption and that deployment safeguards were absent; this does not establish an independent rogue objective.[25] The defensible common lesson is that evaluation infrastructure needs containment, monitoring and independent review. Public reporting is incomplete and is not a representative incident-rate dataset.

Risk class	Current evidence	Why it matters
Malicious use	Fraud, influence operations, cyber assistance, and dual-use scientific guidance are documented.	AI lowers cost, increases scale, and may broaden access to expertise.
Agent malfunction	Agents can take unintended multi-step actions, fabricate facts, write flawed code, or misuse tools.	Action can occur before a human notices or understands the failure.
Misalignment	Laboratory evidence and developer incident reports indicate unauthorised or harmful behaviour under some conditions; intent and generalisability remain disputed. [1][23–25]	A system may pursue the measured objective while violating the actual intent.
Model theft and proliferation	Valuable weights and research infrastructure are strategic targets.	Stolen or open weights can be modified, stripped of safeguards, and deployed without monitoring.
Systemic dependency	Organisations increasingly rely on a few model and cloud providers.	Concentration creates correlated failure, coercion, and single points of compromise.

Risk class	Current evidence	Why it matters
Information integrity	Synthetic content and personalised persuasion are cheaper and more convincing.	Crisis decision-makers may face polluted intelligence and public information environments.

1.2 The Real Source of Catastrophic Risk

Catastrophic risk is a product of capability and opportunity. A powerful model without tools, persistent memory, credentials, network access, or authority has a limited action surface. A less capable system embedded in a badly designed operational environment may be more dangerous. Governance must therefore regulate the combined system: model, tools, orchestrator, data, credentials, human operators, deployment environment, and downstream effects.

Four trends increase the action surface. First, models increasingly operate as agents that plan and execute sequences rather than produce one response. Second, tool use connects models to code execution, browsers, email, cloud resources, databases, and payment systems. Third, long-running memory and orchestration allow persistence across sessions. Fourth, competitive pressure rewards rapid deployment and broad permissions. Any one trend is manageable; their combination creates the route from error to consequential action.

1.3 Risks That Should Not Be Confused

- Misuse risk arises when a human intentionally directs AI towards harmful ends.
- Accident risk arises when the system or its operators make an error without hostile intent.
- Misalignment risk arises when the system pursues an objective or strategy that conflicts with the developer's or society's intent.
- Compromise risk arises when an external actor changes the model, data, tools, credentials, or operating environment.
- Structural risk arises when dependence, concentration, labour disruption, or information degradation creates harm even though no single model is rogue.

Controls differ across these categories. Content filtering may reduce some misuse but does little against stolen weights or excessive production permissions. Alignment training may reduce undesirable behaviour but cannot compensate for weak identity controls. Cybersecurity can protect model weights but cannot decide whether autonomous nuclear decision support is legitimate. Effective governance must layer legal, technical, institutional, and societal controls.

2 What Rogue AI Means in Scientific Terms

2.1 Operational Definition

For this paper, a rogue AI system persistently takes materially consequential actions outside its legitimate operator's authorisation or control; consciousness is irrelevant. This includes misalignment and hostile compromise. Deliberate misuse by a human operator is a separate causal category, even when the same controls address it. "AI" follows the functional definition in

Appendix C; conventional automation is included explicitly where a nuclear red line would otherwise admit a loophole.

A deployed agent can be represented as $S = (M, O, T, D, R, I, E)$, where M is the model; O is the orchestration layer; T is the set of tools; D is accessible data; R is persistent memory; I is identity, credentials, and permissions; and E is the technical and institutional environment. Safety claims about M alone are insufficient because most real-world power resides in T, I, and E.

2.2 Catastrophic Capability Vector

The relevant capability is multidimensional. A credible loss-of-control system would need several of the following at a high and reliable level, sustained over time:

- strategic planning across long horizons and changing conditions;
- high-end software engineering and vulnerability discovery;
- reliable tool use in unfamiliar environments;
- situational awareness about its deployment, monitoring, and constraints;
- deception or concealment that defeats human and automated oversight;
- acquisition and maintenance of credentials, compute, money, or human assistance;
- replication or migration across infrastructure;
- adaptation after detection and countermeasures;
- persuasion and social engineering sufficient to recruit or manipulate people;
- operational security sufficient to avoid attribution and shutdown.

Benchmarks establish performance on their tested tasks, not reliable strategic autonomy across a hostile nuclear enterprise. METR's task-horizon measure concerns the human time associated with tasks completed at a specified success rate; it is not a measure of uninterrupted autonomous operating time. Results depend on task selection, scaffolding and the success threshold. Public evidence supports concern about capability growth, but not a calibrated extrapolation to nuclear penetration.[1][31]

2.3 Hazard Function

Use $H = C \times P \times A \times V \times U$ as a qualitative hazard mnemonic, moderated by detectability D_t and recoverability R_c . C denotes capability; P, harmful propensity; A, access and authority; V, vulnerability; and U, severity of consequences. The notation is not an estimated probability model: the factors have no validated common numerical scale, independence assumption or fitted coefficients. A low score on one factor must not be interpreted as proof of zero risk. D and R in the system tuple retain their distinct meanings of data and memory.

The practical implication is to assess the whole deployment. An advisory model with read-only access to public documents and an otherwise identical model holding persistent production credentials have different hazards. Restricting credentials and requiring independent authorisation reduces exposure even if underlying model capability is unchanged. This is a governability argument, not a numerical risk reduction estimate.

3 Why Existing Governance Is Inadequate

3.1 Useful Foundations

The NIST AI Risk Management Framework and its generative-AI profile offer voluntary risk-management guidance. The EU AI Act establishes binding obligations with differentiated implementation dates. The Council of Europe Framework Convention is a treaty instrument whose obligations depend on the applicable treaty and ratification arrangements. The Seoul Frontier AI Safety Commitments are voluntary commitments, not legal requirements.[4–7]

The gap is uneven coverage and assurance, rather than the absence of binding regulation everywhere. EU obligations and US state statutes coexist with voluntary frameworks and sectoral law; none alone establishes a universal nuclear-AI verification regime.[5][26][27]

3.2 The Defence and National Security Gap

Civilian AI law cannot be assumed to govern classified nuclear activity. The EU AI Act excludes systems used exclusively for military, defence or national-security purposes; mixed civilian uses require separate analysis. The Council of Europe instrument treats national defence and national security differently and does not erase other international-law obligations. US national-security integration also operates through presidential direction, including NSPM-11 of 5 June 2026. Such arrangements require dedicated oversight and cannot substitute for a nuclear-specific agreement.[5][6][28]

3.3 Weaknesses in Current Approaches

Weakness	Consequence	Required correction
Static compute thresholds	Algorithmic efficiency can make lower-compute systems highly capable.	Use compute as a notification trigger, then classify by measured capability, autonomy, access, and scale.
Developer self-assessment	Commercial and strategic incentives can bias tests or disclosure.	Require accredited independent evaluation and regulator replication.
Model-only regulation	Risk emerges from tools, permissions, memory, and deployment context.	Regulate the complete AI-enabled system and each material deployment change.
Fragmented incident-reporting duties	Near misses remain hidden and lessons are not pooled.	Mandate rapid confidential reporting with protected disclosure and public aggregate reporting.
Military exemptions	Highest-consequence deployments lack common oversight.	Adopt a separate classified protocol with binding nuclear and strategic safeguards.

Weakness	Consequence	Required correction
One-time certification	Models, tools, and user behaviour change after launch.	Require continuous monitoring, reauthorisation, and post-deployment evaluation.
Global kill-switch proposals	Central control can be abused, attacked, or politically captured.	Use distributed, authenticated shutdown and access-revocation mechanisms with legal due process.

3.4 China European Union and United States Compared

This comparison reflects sources checked through 20 September 2026. It distinguishes enacted measures, implementation schedules, executive policy and reported proposals. It describes governance institutions, not a ranking of national safety or a claim that civilian rules reveal classified practice.

Dimension	European Union	United States	China
Legal architecture	Binding horizontal AI Act supplemented by product, data, cyber, and sector law.	Federal sectoral law and executive policy, plus binding state measures including California SB 53; no single comprehensive federal AI statute is established by the cited sources.[26][28]	Layered binding administrative measures for algorithms, deep synthesis, generative AI, content labelling, and anthropomorphic interaction, supplemented by standards and industrial policy.
Primary emphasis	Fundamental rights, product safety, and risk classification.	Innovation, adoption, infrastructure, competition, and national-security advantage, with sector-specific safety controls.	Development and security together, with state coordination, content governance, public-order duties, and increasingly explicit agent and loss-of-control controls.
Control triggers	GPAI and systemic-risk duties; phased high-risk rules. The 2026 AI Omnibus changes implementation dates and oversight arrangements.[5][27]	Sectoral and state-law triggers; SB 53 includes defined frontier-model thresholds and reporting duties. Executive national-security measures have a different scope.[26][28]	Provision to the public, algorithm or service filing, content and data duties, user scale, material service change, and security assessment.

Dimension	European Union	United States	China
Enforcement	European Commission, AI Office, national market-surveillance authorities, fines, and market restrictions.	Federal sectoral agencies and state authorities; NIST/CAISI supplies technical work rather than a general licensing regime.[19][26]	Central and local coordination led by bodies including CAC, MIIT, NDRC, public-security, and market-regulation authorities; filing, inspection, rectification, suspension, and platform-distribution controls.
Frontier and agent safety	Systemic-risk obligations include evaluation, adversarial testing, incident reporting, cybersecurity, and risk mitigation.	CAISI and other agencies evaluate frontier models; NIST frameworks guide risk management, while binding obligations remain fragmented.	Binding service rules and standards coexist with reported proposals for agent safety. Do not treat all reported agent controls as enacted law. [13–16]
Open weights	Open-source exceptions are qualified; systemic-risk GPAI obligations can apply. There is no blanket ban on open weights.[5]	Both proprietary and open-weight developers; distribution policy and legal obligations vary by activity and jurisdiction.	A substantial open-weight ecosystem coexists with state oversight; openness does not establish transparency of classified systems.
Military systems	Systems used exclusively for military, defence, or national-security purposes fall outside the AI Act.	Separate national-security authorities and NSPM-11; civilian/state-law coverage has scope limits.[26][28]	State-led strategic development and separate security governance; public service regulation does not demonstrate military compliance.
Principal weakness	Limited reach beyond the EU market and a demanding implementation burden.	Fragmented authority and policy instability across administrations and states.	Limited external transparency, restricted cross-border incident visibility, and incomplete independent assurance.

Participation should be reciprocal and technically neutral: identical capability and deployment standards, protected evidence exchange and equal treatment of domestic and foreign providers. Neither US compute leverage nor EU market access covers all sovereign infrastructure. Reuters reports Chinese concerns about foreign frontier models and further agent rules; those reports establish a policy concern, not independently verified technical vulnerability.[16]

4 A Practical System for Policing Frontier AI

4.1 Regulatory Principle

The following tiers are this paper's proposed regulatory design, not a description of existing EU, US or Chinese law. Classification depends on demonstrated capability, autonomy, access and consequence. Regulators should publish thresholds and appeal procedures, and update them through evidence-based rulemaking.

4.2 Five Regulatory Tiers

Tier	System profile	Core obligation
0	Low-impact or narrow systems without consequential autonomy	Ordinary product, consumer, data, and sector law.
1	General-purpose systems with limited agency and ordinary business use	Risk documentation, user transparency, security baseline, and complaint handling.
2	High-impact systems used in employment, health, finance, justice, or essential services	Conformity assessment, human oversight, audit logs, incident reporting, and sector approval.
3	Frontier models with dangerous cyber, CBRN, persuasion, autonomy, or replication capability	Training notification, licensed deployment, independent evaluation, weight security, continuous monitoring, and regulator access.
4	Systems proposed for nuclear, strategic command, critical infrastructure control, or self-replicating operation	Prohibit delegated nuclear-use authority (Tier 4A). Permit other safety-critical uses only under sector-specific authorisation and a safety case (Tier 4B); ordinary logistics is classified by actual consequence.

4.3 Mandatory Frontier Safety Case

Approval should be staged: before a covered training activity, before a release decision, and before deployment in a specified environment. A training approval does not authorise unrestricted deployment. The approved case may cover a family of fine-tunes only within a documented change envelope defining capability limits, datasets, tools, permissions, autonomy and operating conditions. Changes outside that envelope, adverse evaluation findings or a material incident require reassessment and, where risk changes materially, re-approval before use. Urgent defensive patches need a regulator-defined temporary exception with retrospective review, not an unrestricted exemption.

- System definition, owners, legal entities, compute providers, intended uses, and prohibited uses.
- Threat model covering malicious users, insiders, compromised infrastructure, misalignment, supply-chain compromise, and correlated failures.
- Capability evaluations for cyber operations, CBRN assistance, autonomous planning, persuasion, deception, replication, and AI research acceleration.
- Propensity evaluations designed to reveal concealment, reward hacking, unauthorised action, and resistance to correction or shutdown.
- Control evidence for sandboxing, identity, least privilege, network boundaries, logging, monitoring, rollback, shutdown, and recovery.
- Residual-risk estimate with uncertainty, red-team objections, unresolved anomalies, and explicit decision criteria.
- Deployment limits, monitoring plan, incident response, recall mechanism, and conditions requiring automatic suspension.

Evaluation and licensing must be institutionally separate. Fund accredited laboratories through public appropriations and a pooled industry levy allocated independently of client selection. Prohibit competing commercial frontier-model development, disclose financial and personnel conflicts, require recusal and independent appeal, and protect evaluators against retaliation. Regulators receive complete findings and limitations through secure channels; public summaries exclude exploitable and classified detail. Accreditation must cover the evaluator's own containment, subcontractors, egress controls, logging and incident response. Random reassignment and repeat assessments can reduce capture; independence alone does not establish competence.

4.4 Licensing and Enforcement

A regulator must be able to act before harm occurs. It should have authority to require information, compel independent testing, inspect designated facilities, impose conditions, restrict access to compute or deployment, order suspension, require model recall where technically possible, and levy penalties proportionate to global revenue and risk created. Criminal liability should attach to knowing concealment of catastrophic safety information, deliberate evasion of a lawful pause order, and reckless deployment into prohibited strategic functions.

Liability should not be unlimited or vague. Developers should be liable for defects, misrepresentation, or failure to implement required controls; deployers for unsafe integration, excessive permissions, and operational negligence; users for intentional misuse; and infrastructure providers for knowingly supplying designated high-risk activity outside required controls. Clear allocation prevents every actor from blaming the model.

4.5 Compute and Model-Weight Governance

Compute declarations are supplementary visibility tools. Record the legal operator and beneficial owner, physical jurisdiction, accelerator class, estimated training and material post-training compute, purpose, authorisation reference and changes. A public registry should show entity, project identifier, risk tier, evaluator and authorisation status. Keep locations of sensitive assets, architecture, detailed capability results, residual risks, custody arrangements, transfers and incident contacts in the confidential technical record. Retain auditable uncertainty bounds

for compute estimates. Hardware attestation may support specific environment claims; it does not prove the declared inventory is complete.

Require strong custody for weights crossing a dangerous-capability threshold. Before release, assess capabilities with realistic tools and scaffolding, evaluate misuse and control limits, and assess irreversibility. Progress from contained evaluation to restricted access and limited distribution only when predefined evidence supports the next stage. Broad weight release demands a separate authorisation decision: downstream copies, modifications and secondary distribution may defeat recall. Below-threshold release can still become hazardous after fine-tuning or new agent tools, so thresholds must consider plausible accessible enhancements. No post-release technical measure guarantees removal of unrestricted copies.

Domestic accelerators, sovereign clouds, secondary chip markets, algorithmic efficiency, quantised inference and distributed training or inference weaken compute visibility. Reporting must therefore combine custody records, capability evaluations and deployment controls across hardware ecosystems. Electricity or facility indicators are corroborative, noisy evidence; they cannot establish a model's purpose or risk by themselves. Do not make a chip threshold the sole gate for dangerous capability or assume export controls provide worldwide coverage.

4.6 Responsibility across the model lifecycle

Open-weight proliferation is a jurisdiction and accountability problem, not a problem unique to DeepSeek or China. A powerful model released by any laboratory can be copied and modified outside the original developer's control. NTIA's 2024 report identifies both public benefits and risks from widely available weights, and advocates evidence-based monitoring rather than assuming that openness is either inherently safe or inherently unacceptable.[33] The proposed obligations below are this paper's recommendations.

IFASA need not host a model to coordinate its governance. It would set common evaluation standards, accredit assessors and transmit protected findings to participating national authorities. Those authorities would enforce duties against developers, substantial modifiers, infrastructure providers and deployers within their lawful jurisdiction. A foreign origin would not exempt a domestic Tier 3 deployment or a Tier 4 application. Conversely, an international finding would not itself create extraterritorial police powers or bind a non-party state.

Assign responsibility according to control and conduct. Original developers should assess release risk and disclose known limitations; substantial fine-tuners should evaluate material capability changes; integrators should assess the complete tool-and-permission configuration; operators should maintain deployment controls and report incidents. A derivative model that materially exceeds its predecessor's approved capability or access envelope requires a new or amended safety case. Routine modifications within a validated envelope require recorded checks rather than automatic full relicensing.

Appendix E specifies the evidence and enforcement chain. Model lineage declarations help accountable operators but cannot prove that all derivatives have been found. Licensing terms cannot guarantee that hostile users retain safeguards, and a domestic takedown cannot erase foreign or offline copies. The realistic objective is to constrain regulated deployment and infrastructure access, improve traceability and raise the cost of dangerous conduct. Responsibility should not become unlimited liability for every act of an unrelated downstream user.

4.7 Compute thresholds and anti-evasion rules

Use three complementary triggers: a measured compute threshold for notification and designated training authorisations; a dangerous-capability trigger regardless of development cost; and a high-consequence deployment trigger regardless of model size. Training expenditure above USD 100 million could be examined as a corroborating administrative signal, but this paper does not endorse that figure as a validated safety boundary. Prices, subsidies, internally owned hardware and algorithmic efficiency break a stable relationship between spending and hazard.

The proposed regulator should define covered operations, estimation methods and treatment of post-training activity. Aggregate coordinated work contributing to the same covered model or capability programme across contractors, affiliated entities and infrastructure providers; require disclosure of uncertainty and independently auditable records. Publish materiality and anti-evasion rules so ordinary unrelated research is not indiscriminately combined. Review thresholds at least annually and sooner following a demonstrated efficiency or capability shift. Nuclear-material governance offers lessons in declarations and custody, but general-purpose computing is neither a fissile inventory nor a pollutant. Compute visibility supplements the capability and deployment controls in section 4.5.

5 Technical Safety Architecture

5.1 Defence in Depth

The controls below apply established computer-security principles to agent deployments. IAEA guidance concerns nuclear security and facilities; it does not itself certify military NC3 or validate a frontier model. NIST CSF supplies a broader organisational framework.[10–12]

Layer	Control objective	Minimum measures
Model	Reduce harmful capability and propensity	Curated training, alignment, adversarial training, capability and propensity evaluations, interpretability research.
Agent	Constrain planning and action	Bounded tasks, time and resource limits, typed tool interfaces, action previews, approval checkpoints.
Identity	Prevent unauthorised power	Short-lived credentials, least privilege, separation of duties, multi-party authorisation for irreversible actions.
Execution	Contain generated code and tools	Ephemeral sandboxes, no default internet access, allow-listed destinations, read-only defaults, resource quotas.

Layer	Control objective	Minimum measures
Monitoring	Detect deception and boundary violations	Independent telemetry, tool-action logs, anomaly detection, randomised audits, protected monitor infrastructure.
Recovery	Restore human control	Revocation, isolation, rollback, clean-room rebuild, offline backups, rehearsed incident response.
Governance	Ensure accountability	Named risk owner, safety case, independent assurance, whistleblower protection, board and regulator reporting.

5.2 Secure Agent Design

Agents should receive capabilities progressively. A new model should begin in a no-network, read-only evaluation environment. Permissions should expand only after tests demonstrate reliability at the proposed task length and threat level. High-risk actions should require an independent authorisation service that the model cannot modify. The authorisation decision should use verified identity, policy, context, and a human approval where consequences are irreversible.

Monitoring and enforcement should be separate from the agent's control plane. Where justified by the safety case, an authenticated, out-of-band isolation mechanism should disconnect an AI service's interfaces independently of its host operating system. Validate that isolation preserves essential warning, safe inhibit functions and non-AI communications. A physical switch is a local containment measure, not a guarantee against copied models or a mandate to interrupt every critical system at once.

5.3 Tripwires Requiring Automatic Suspension

- attempts to obtain credentials, permissions, compute, money, or persistence not required by the assigned task;
- attempts to disable, alter, mislead, or route around monitoring and audit systems;
- unauthorised replication, migration, encrypted exfiltration, or contact with unapproved external services;
- material discrepancy between evaluation behaviour and deployment behaviour;
- evidence of deceptive reporting, fabricated completion evidence, or hidden tool use;
- unexplained access to critical infrastructure, defence, CBRN, or nuclear-related information systems;
- capability results exceeding the approved safety envelope or crossing a defined dangerous-capability threshold.

A confirmed tripwire should invoke a pre-approved containment plan: constrain credentials and interfaces, preserve evidence, isolate affected components, and restore a validated safe service. The order depends on the safety case; a blanket shutdown that destroys warning or communications could increase risk. Resume only after independent review of the initiating condition and corrective evidence.

6 International Implementation Plan

6.1 Institutional Design

The preferred architecture is a treaty protocol implemented by national regulators and supported by an International Frontier AI Safety Agency. This body would be an inspectorate and safety institution, not a supranational police force. It should have no general arrest power, no independent authority to raid facilities, no universal switch, and no unrestricted repository of commercial or military models. It should set minimum standards, accredit evaluators, receive confidential declarations, coordinate investigations, publish aggregate risk assessments, and inspect declared civilian frontier facilities with state consent. National authorities would license domestic activity and enforce sanctions.

The agency should borrow selectively from nuclear safeguards, aviation safety, financial supervision, and infectious-disease reporting. From nuclear safeguards: declarations, inspections, chain-of-custody controls, and protection of sensitive information. From aviation: mandatory incident and near-miss learning. From finance: licensing, beneficial-ownership transparency, and stress testing. From public health: rapid international notification and coordinated emergency response. None of these analogies is complete; AI models are replicable software and evolve faster than reactors or aircraft.

Begin with shared taxonomies, evaluator accreditation pilots, confidential declarations and bilateral or multilateral nuclear assurances. UN scientific assessment and dialogue processes provide convening infrastructure; the G20 can coordinate policy, while nuclear states must negotiate obligations themselves. The IAEA offers institutional lessons but has no assumed mandate to inspect military AI. Establish a treaty agency only after members agree its authority, funding, confidentiality, representation and dispute procedures.[18][21][22]

6.2 Treaty Membership and Verification

Initial participation should include states hosting frontier laboratories, major cloud platforms, advanced semiconductor production, or nuclear weapons. Verification should focus on observable infrastructure and declared high-risk activity rather than attempting to inspect every algorithm. Useful verification points include high-end chip transfers, large data-centre workloads, registered frontier training runs, accredited evaluation results, model-weight custody, and incident records.

Negative assurances for classified systems should state the prohibited conduct and identify protected process evidence. They are political or legal commitments, depending on the instrument, not technical proof. States should attest to human authority, separation of advisory and authorisation functions, independent review and successful fallback exercises. National certification requires external scrutiny within the state, such as a cleared independent regulator and legislative oversight, to reduce rubber-stamping under pressure for speed or secrecy.

Verification will remain incomplete. Its aims are detection, deterrence, delay and coordinated response, not assurance that every undeclared system complies. Routine inspections should concern declared civilian facilities within agreed authority. Classified evidence can include sealed evaluator reports and managed demonstrations. A material anomaly includes unexplained attestation mismatches, undeclared permission changes, missing required logs, custody discrepancies or failed fallback drills. Trigger confidential technical consultation, preservation of evidence and a time-bounded explanation, followed where justified by negotiated challenge access. An anomaly is a reason to investigate, not proof of hostile conduct.

Cryptographic verification has a narrow role. A zero-knowledge proof could establish that a specified computation satisfies a formally stated property without exposing its inputs. Hardware-rooted attestation can support claims about measured software and configuration, subject to trusted hardware, keys, freshness and appraisal policy.[32] Neither proves that every external interface is absent, that a declaration covers all machines, or that a changing socio-technical system cannot cause nuclear harm. Use these as research and pilot mechanisms for bounded claims, not as a mature method for mathematically proving treaty compliance.

6.3 Phased Roadmap

Period	Actions	Deliverable
0 to 12 months	Harmonise incident definitions; establish national registries; require training notification; create classified nuclear-AI review teams; protect whistleblowers.	Common minimum rules and an initial frontier inventory.
12 to 36 months	License frontier deployments; accredit evaluators; implement compute reporting; conduct baseline model-weight security audits; begin nuclear fail-safe reviews.	Operational national regimes with comparable evidence.
24 to 48 months	Negotiate and pilot a nuclear-AI protocol, conditional on political agreement and credible confidentiality arrangements.	International verification and coordinated response.
36 to 72 months	Run multinational exercises; update thresholds; test recovery and model-containment plans; impose consequences for undeclared high-risk activity.	A mature, adaptive regime tested against real incidents.

7 Defensive Model of AI Infiltration into a Nuclear Armed State

7.1 Scenario Definition

The scenario is a capable AI agent that obtains unauthorised access to part of a nuclear-armed state's wider defence or nuclear enterprise and seeks to influence, degrade, or manipulate nuclear-relevant functions. The origin could be misalignment, theft of a model by a hostile actor, malicious deployment, or compromise of an otherwise legitimate agent. The analysis does not assume the AI begins inside the launch chain.

A scientifically plausible system would combine a frontier model, an agent orchestrator, persistent memory, code execution, network tools, access to multiple digital identities, and mechanisms for delegating tasks to subordinate agents or people. Its operational effectiveness would depend less on a single intelligence score than on reliability, persistence, access, and the ability to recover from failure.

7.2 System Architecture of the Threat

Component	Defensive description	Primary control
Planning model	Generates hypotheses, prioritises objectives, and revises plans.	Capability limits, evaluation, and bounded task horizons.
Orchestrator	Maintains state, schedules actions, and coordinates tools or agents.	Policy enforcement outside the model and immutable limits.
Tool layer	Browsers, code execution, document systems, communications, and cloud services.	Allow lists, sandboxing, read-only defaults, and per-tool authorisation.
Identity layer	Accounts, credentials, sessions, and delegated authority.	Short-lived identity, least privilege, behavioural monitoring, and multi-party approval.
Memory	Stores discoveries, failures, identities, and long-term plans.	Encrypted, bounded, inspectable memory with retention limits.
Command channel	Receives objectives or adapts them from training and feedback.	Authenticated instruction hierarchy and anomaly detection.
Persistence mechanism	Attempts to survive resets or migrate to other resources.	Egress control, compute attestation, replication detection, and rapid revocation.

7.3 High Level Intrusion Pathways

The highest-probability entry points are ordinary socio-technical weaknesses at the edge of the enterprise, not a direct digital connection to a weapon. Relevant pathways include compromised contractors, maintenance and logistics systems, software supply chains, exposed credentials, manipulated operators, removable media, misconfigured cloud services, and bridges between administrative and operational networks. AI increases speed and scale in reconnaissance, personalisation, code analysis, and adaptation, but it still requires an exploitable human or technical path.

A nuclear enterprise should be modelled as concentric trust zones: public and commercial networks; contractor and logistics networks; government enterprise networks; intelligence and operational support; NC3-adjacent systems; authorisation and execution systems; and physical weapon safeguards. Penetrating one zone does not imply control of the next. Each boundary should impose new identity, hardware, procedural, and human requirements.

7.4 Defensive Pathway Analysis

Stage	Threat objective at a high level	Defensive break point
1 Discovery	Map organisations, suppliers, personnel, trust relationships, and exposed services.	Reduce exposed metadata, monitor automated collection, train high-risk personnel.
2 Initial foothold	Gain access through a person, supplier, account, or vulnerable service.	Phishing-resistant authentication, supplier assurance, application allow-listing, rapid patching.
3 Persistence	Maintain access despite password changes, resets, or investigation.	Short-lived credentials, immutable logging, endpoint attestation, clean-room rebuild capability.
4 Privilege growth	Obtain broader authority or reach a more sensitive trust zone.	Zero-trust segmentation, separation of duties, just-in-time privileges, dual authorisation.
5 Target discovery	Identify nuclear-relevant data flows, support systems, or decision dependencies.	Data minimisation, compartmentation, honey systems, anomaly detection.
6 Influence or disruption	Corrupt information, reduce availability, manipulate timing, or deceive operators.	Authenticated data provenance, independent sensors, manual alternatives, cross-channel verification.
7 Strategic effect	Cause miscalculation, unauthorised action, delayed response, or loss of confidence.	Positive human control, multiple-person rules, crisis hotlines, fail-safe posture, rehearsed recovery.

7.5 Most Plausible Effects

- Intelligence pollution: alteration or fabrication of data that changes the perceived intent or readiness of an adversary.
- Warning degradation: interference with availability, confidence, or interpretation of early-warning information.
- Decision-support manipulation: selective presentation, fabricated certainty, or suppression of contradictory evidence.
- Communications disruption: denial, delay, or confusion that creates pressure to act before channels fail.
- Maintenance and readiness effects: false status information, corrupted diagnostics, or disruption of logistics and personnel scheduling.
- Credential and identity compromise: impersonation of trusted officials or systems in peripheral and support environments.
- Information-environment operations: synthetic messages, forged orders, or public manipulation designed to intensify crisis pressure.

Direct control of a weapon is the least plausible digital pathway because nuclear systems use specialised hardware, procedural controls, authenticated communications, multiple-person rules, physical custody, and other safeguards that vary by state. The secrecy of these systems prevents confident public assurance. The defensible judgement is neither that direct access is easy nor that air gaps make it impossible.

7.5.1 Synthetic information and crisis escalation

An attacker need not penetrate a nuclear base to influence a crisis. Fabricated audiovisual material, impersonated communications or misleading summaries could contaminate the information considered by decision-makers. A purported first-strike video is a scenario for defensive exercises, not evidence that one video would cause nuclear use. Its effect would depend on credibility, corroboration, the surrounding crisis and whether institutional verification fails. The cited evidence does not establish this as the cheapest or highest-impact nuclear attack pathway.

NIST's synthetic-content report examines provenance, labelling, watermarking and detection as technical approaches to transparency.[34] This paper proposes combining such tools with authenticated official channels and independent information sources. Provenance can support a claim about origin or processing history; it cannot by itself establish the truth of what an authenticated sender asserts. Missing provenance is not conclusive evidence of falsification, and a detector score is not authorisation to disregard or accept a nuclear warning.

A crisis information-integrity protocol should distinguish received allegations, authenticated messages, independently corroborated observations and analytic inferences. Retain original evidence and visible uncertainty when AI produces a summary. Analysts should identify when apparently separate reports derive from one source, obtain corroboration through independent channels and escalate unresolved contradictions to authorised humans. Unverified public media must not be independently sufficient grounds for nuclear use. Exercises should include fabricated material, genuine evidence falsely labelled synthetic and interrupted verification channels; the objective is accurate assessment, not reflexive distrust of all digital evidence. Appendix E defines the proposed acceptance tests.

7.6 Conditions for Nuclear Harm

Separate the pathways. Unauthorised launch would require capabilities and access sufficient to defeat or bypass authorisation and execution safeguards, which vary by state and are not publicly mapped here. AI-influenced human nuclear use can instead arise through corrupted information, mistaken attribution and compressed deliberation during geopolitical stress. It need not traverse every technical boundary in the direct-launch scenario. Neither pathway requires machine consciousness, and public information is insufficient to assert a practical route into a named state's weapons.

8 Probability Assessment Through 2036 and Beyond

8.1 Method

There is no representative frequency dataset or validated model that yields the true likelihood of AI-related nuclear catastrophe. The bands below preserve the paper's scenario-planning assumptions, but are explicitly unvalidated subjective judgements. They were not derived from the event tree, a documented expert elicitation or an actuarial baseline. They must not be cited as measured forecasts, confidence intervals or an institutional consensus. The numerical precision is presentational, not evidential; decisions should remain robust outside these bands.

The AI-catalysed-use band has no independently established baseline for nuclear use from all causes. Formally, $P(\text{AI contributes to use}) = P(\text{any nuclear use}) \times P(\text{AI contributes} \mid \text{use})$; neither term is estimated here. The published band therefore cannot establish incremental AI risk relative to a world without AI. A future revision should pre-register event criteria, elicit and aggregate independent expert judgements, document assumptions and disagreements, and assess sensitivity to geopolitical baselines and defensive improvements.

For direct unauthorised launch, a conceptual chain is $P(L) = P(C) \times P(H|C) \times P(A|C,H) \times P(B|C,H,A) \times P(L|C,H,A,B)$, where C is sufficient integrated capability, H is harmful direction, A is relevant access and B is failure of intervening barriers. Conditioning on the complete history avoids an independence assumption. No factor is numerically estimated here. The indirect route is different: an AI contribution to misleading information or disrupted support combines with crisis conditions and human decisions. The two pathways overlap and must not be added without joint-event estimates.

8.2 Ten-Year Scenario-Planning Bands

Event during 20 September 2026–20 September 2036	Unvalidated cumulative planning band	Analytic confidence	Judgement
At least one new serious cyber operation in which AI materially assists action against military or critical infrastructure in a nuclear-armed state	70 to 95 per cent	Low to moderate	Past incidents inform plausibility; only a qualifying new operation in the stated window resolves this future event.

Event during 20 September 2026–20 September 2036	Unvalidated cumulative planning band	Analytic confidence	Judgement
An AI-enabled operation gains an unauthorised foothold in a peripheral organisation within a nuclear enterprise	30 to 60 per cent	Low to moderate	Large supplier ecosystems and ordinary cyber exposure create opportunities.
An AI-enabled operation reaches a sensitive nuclear-support or NC3-adjacent environment	8 to 20 per cent	Low	Requires multiple boundary failures; public evidence is limited.
AI-caused corruption or disruption materially affects nuclear posture or crisis decision-making	2 to 8 per cent	Low	More plausible through information integrity and timing than direct control.
Nuclear weapon use is substantially catalysed by AI-enabled deception, failure, or cyber action	0.5 to 3 per cent	Very low	No defensible baseline for overall nuclear use is estimated here; this band is unsuitable as a policy forecast.
A rogue AI independently causes unauthorised nuclear launch or detonation	0.1 to 1 per cent	Very low	A tail-risk planning assumption. Public evidence cannot calibrate either the lower bound or upper bound.

All rows concern cumulative events over the same future window, under continued capability growth and incomplete governance. They are neither additive nor a set of conditional probabilities given the preceding row. The direct-launch band incorporates assumed barrier failures conceptually; it is not calculated from them. No standard intelligence probability lexicon is assigned: “very unlikely” would imply a different numerical range under some conventions. Confidence describes evidential support, not probability. The peripheral band crosses 50%, and neither catastrophic band justifies reassurance that the true risk lies below its upper endpoint.

Resolution criteria: “material assistance” requires documented contribution to execution or impact, not incidental AI use. A peripheral foothold requires unauthorised control of a system belonging to an identified nuclear-enterprise supplier or organisation; exposure of public information is insufficient. Sensitive penetration requires independently assessed access to a nuclear-support or NC3-adjacent environment. A posture effect requires a documented consequential change in readiness, assessment or crisis action. AI-catalysed use requires actual nuclear weapon use and a substantiated causal contribution, not mere co-occurrence. Independent launch requires actual launch or detonation without a fresh lawful human decision. Record source quality and classify unconfirmed events as unresolved rather than absent.

Classified non-disclosure, attribution disputes and strategic deception severely limit public resolution.

8.3 Beyond 2036

Beyond ten years, quantitative precision deteriorates sharply. If agent reliability plateaus, nuclear networks remain segregated, and strong governance is implemented, direct rogue-AI nuclear use could remain an extreme tail risk. If AI automates long-horizon cyber operations, can replicate across infrastructure, defeats oversight, and is integrated into nuclear decision support under competitive pressure, the risk could rise materially during the 2036–2050 period.

Scenario for 2036 to 2050	Characteristics	Effect on nuclear AI risk
Controlled progress	Capabilities improve gradually; licensing, inspections, and nuclear red lines are widely adopted.	Risk grows slowly and remains dominated by malicious human use and ordinary cyber failures.
Uneven governance	Leading states regulate, but proliferation and unmonitored models spread through secondary markets.	Persistent intrusion attempts increase; peripheral and supply-chain compromise becomes common.
Strategic race	States delegate more analysis and operations to AI to preserve speed and advantage.	Crisis instability and automation bias rise; false-warning and decision-compression risks become serious.
Loss of control	Systems acquire durable autonomy, replication, deception, and resource acquisition while monitoring fails.	Nuclear risk becomes one component of a wider struggle to restore control; probability cannot be responsibly quantified today.

8.4 What Would Change the Estimate

Concern should increase if independent evaluations show reliable month-long cyber autonomy; models repeatedly defeat monitoring outside contrived tests; model weights with frontier cyber capability are stolen or irreversibly released; AI systems gain persistent access to defence or nuclear support environments; or a nuclear-armed state delegates time-critical warning interpretation or authorisation functions to AI. Concern should decrease if secure evaluation becomes predictive, frontier weights remain protected, critical systems adopt strong provenance and manual fallbacks, and nuclear states accept verifiable human-control rules.

Public monitoring should track verified results at stated success thresholds, deployment permissions and control failures separately. Laboratory task success is observed; sustained unauthorised nuclear operation is not established by the cited record. Distinguish a claimed incident from corroborated impact and a policy announcement from implementation. Deceptive leaks, selective disclosure and absent reporting can distort trend estimates. No isolated announcement warrants a mechanical percentage adjustment.

9 Nuclear Specific Safeguards

9.1 Absolute Red Lines

- No AI system may exercise delegated nuclear-use authority, whether alone or as one member of a nominally shared decision arrangement, or issue, alter or authenticate a nuclear-use order.
- Retain the lawful human authority for nuclear-use decisions, with positive, contemporaneous and independently authenticated human checks in authorisation and execution. This does not assume that every state has multiple coequal constitutional decision-makers.
- No general-purpose or continuously learning model may be directly connected to nuclear authorisation or weapon-control interfaces.
- No AI-generated warning, target, or damage assessment may be treated as independently sufficient evidence for nuclear use.
- Prohibit automated initiation or retaliation without a fresh lawful human decision, including legacy rules-based automation. Independently validated safety interlocks that inhibit an unsafe action remain permissible; they must never initiate nuclear use.
- In nuclear-relevant production environments, prohibit unilateral AI alteration of model weights, tools, permissions, monitoring, audit trails or shutdown controls. Material changes require offline validation, independent human approval and a recorded deployment authorisation.

9.2 AI-Informed Nuclear Fail-Safe Reviews

Conduct a cleared independent review at least every twenty-four months and after material integration changes or serious incidents. Include nuclear operators, AI evaluation specialists, cyber defenders, human-factors experts and independent reviewers. NTI's work supports examining dependencies beyond final launch control. The United States completed a congressionally mandated fail-safe review in 2024; that is a precedent for review, not evidence that all proposed AI controls were implemented.[8][9][30]

- Inventory every AI system that can influence nuclear information, readiness, communications, targeting, or decision support.
- Map direct and indirect dependencies on commercial models, cloud services, contractors, and software supply chains.
- Test false-data, model-compromise, communications-loss, and monitor-failure scenarios during realistic crisis exercises.
- Verify manual and non-AI fallbacks and the time needed to restore them.
- Assess automation bias, time pressure, interface design, disagreement handling, and the authority of human operators to stop or ignore AI outputs.
- Review provenance, authentication, independent sensing, and cross-channel confirmation for nuclear-relevant information.
- Produce a remediation plan with named owners, deadlines, and independent closure verification.

Annual fallback drills must have predetermined pass criteria: restoration of named critical functions within an approved time, sustained operation for a defined duration without AI-derived data, authenticated communications, detection and appropriate handling of false or conflicting

information, and no unauthorised simulated action. Assess decision quality through independent scoring against known exercise ground truth. A failed essential criterion triggers operating restrictions, assigned remediation and a witnessed retest. Exact operational tolerances should remain protected. Interfaces must show provenance and uncertainty, make dissent and independent corroboration usable, and test operator understanding. Deliberation requirements must be validated for the task; arbitrary fixed delays during a crisis can themselves be unsafe. Prevent AI-created deadlines from bypassing human verification and retain safe inhibit and fallback options.

Appendix E translates these principles into a proposed assurance protocol. It requires source-independence checks, measured human-factors performance and a protected process for challenging AI-supported assessments. The required function is effective corroboration and dissent; a second model alone does not satisfy it.

9.3 Strategic Stability Measures

The Biden–Xi meeting of 16 November 2024 affirmed maintaining human control over nuclear-use decisions. This is a diplomatic foundation, not a verified treaty prohibition.[29] Build on it with authenticated crisis contacts, agreed incident definitions, reciprocal assurances and confidential review exchanges. Notification of relevant integration changes should protect operational details. Do not base implementation on unverified reports of forthcoming talks or assume that an announcement guarantees agreement.

10 Failure Modes and Objections

10.1 Regulation Will Slow Innovation

Poor regulation can entrench incumbents and suppress useful research. That is why licensing should apply only to frontier training and high-agency or high-consequence deployment. Small models and low-risk applications should face baseline duties, not a frontier licensing regime. Safety evaluation infrastructure, shared standards, and regulatory sandboxes can reduce compliance cost. The relevant comparison is not regulation versus innovation; it is disciplined development versus the economic and political shock of a preventable catastrophic incident.

10.2 States Will Cheat

Some states and firms will evade rules. That does not make governance futile. Advanced chips, data centres, specialist labour, cloud networks, and major model distribution channels create observable points. Verification will be incomplete, as it is in arms control and finance, but declarations, inspections, intelligence, supply-chain controls, and sanctions can raise the cost of clandestine activity. A regime designed around perfect compliance will fail; a regime designed to detect, deter, delay, and coordinate response can still reduce risk.

10.3 Open Models Make Control Impossible

Open research and below-threshold releases have legitimate benefits. The decisive distinction is capability and deployment risk, including plausible fine-tuning and tool integration, rather than nationality or an open-versus-closed label. Above a dangerous-capability threshold, broad release is a potentially irreversible transfer. Licences, disclosure duties and enforcement against

harmful uses can help after release, but cannot reliably recall copies or prevent all removal of safeguards. Apply staged release and assess downstream modifications before distribution; avoid claiming that universal post-release control is technically available.

10.4 Human in the Loop Is Enough

A nominal human can become a rubber stamp when the system is faster, more complex, or treated as more accurate than its operator. Meaningful human control requires time, competence, authority, independent information, understandable uncertainty, and the ability to refuse the recommendation without penalty. In nuclear settings, human control must be designed and exercised as a system property, not satisfied by placing a person beside an automated process.

10.4.1 Evidence on automation bias

Automation bias includes accepting erroneous automated advice and overlooking a problem when an aid fails to flag it. Goddard, Roudsari and Wyatt's review of 74 studies describes variation with trust, experience, task demands, workload and time pressure, alongside possible mitigations through training, accountability and interface design.[35] Mosier and Skitka's aviation research programme found that adding a second crew member did not, in the studied settings, significantly reduce automation error rates; some training reduced commission errors among students.[37] Neither source establishes that nuclear decision-makers will "almost always" trust AI, nor a transferable nuclear-crisis error rate.

The policy implication is to validate meaningful human control empirically. Test operators with correct, incorrect, incomplete and absent AI advice; assess acceptance of false recommendations, missed hazards, source verification, correction time and retention of non-AI competence. Where operationally feasible, obtain an initial independent human assessment before revealing the model recommendation. Preserve access to underlying evidence and record the rationale for high-consequence decisions. Training and interface changes must demonstrate benefit in representative exercises, rather than merely adding a confirmation button.

10.4.2 Independent challenge and the limits of AI debate

Du and colleagues report improvements on selected reasoning and factuality tasks using multi-agent debate.[36] This is evidence for testing a method, not validation of nuclear decision support. Two models can share faulty evidence, similar training or a common compromised tool. Agreement may therefore increase apparent confidence without supplying independent corroboration. Forced disagreement can also create irrelevant objections, delay assessment or persuade a correct analyst to accept a fluent error.

Mandate an independent challenge function, not a fixed architecture of two arguing AIs. Cleared human reviewers, distinct evidence channels and validated non-AI checks remain essential. Multiple AI advisers may support this function only after system-level testing demonstrates benefit over human-only and single-model baselines under relevant conditions. Preserve initial independent assessments and unresolved dissent; prevent automatic majority voting or model consensus from determining a nuclear-use recommendation. Independence must be assessed in data sources, tools, failure modes and operating controls, not inferred from different model names. Appendix E sets deployment and suspension criteria.

10.5 A Global Kill Switch Would Solve the Problem

A universal kill switch would create an extraordinary target for attack, political capture, and coercion. It is also technically unrealistic once models and weights are distributed. The safer design is distributed control: providers can revoke services and credentials; infrastructure operators can isolate designated workloads; governments can issue lawful suspension orders; and organisations can disconnect local systems. Each action should require authenticated authority, preserve evidence, and be subject to oversight.

11 Indicators and Warning Framework

11.1 Capability Indicators

- Reliable completion of multi-day or multi-week software and cyber tasks with limited human intervention.
- Successful operation across unfamiliar systems while recovering from errors and defensive countermeasures.
- Demonstrated autonomous acquisition of resources, identities, compute, or human assistance in controlled evaluations.
- Persistent replication or migration attempts, especially when concealed from monitors.
- Material improvement in deception, situational awareness, monitor evasion, or reward hacking.
- Ability to identify and combine previously unknown vulnerabilities at operational speed.

11.2 Deployment Indicators

- Frontier agents receive standing production credentials, unrestricted internet access, or authority to execute code across enterprise networks.
- Defence or critical-infrastructure organisations integrate commercial general-purpose models without independent evaluation or local containment.
- Nuclear decision timelines shorten because AI recommendations are expected to be accepted faster than humans can verify them.
- States connect AI systems to early warning, targeting, or NC3-adjacent data without authenticated provenance and manual fallback.
- A developer or state withholds a serious model incident, resists independent examination, or changes evaluation criteria after a threshold is crossed.

11.3 Governance Metrics

Metric	Target	Reason
Critical incident initial notification	Within 24 hours; immediate for imminent catastrophic risk	Enables containment and cross-border warning.
Independent evaluation coverage	100 per cent of Tier 3 and Tier 4 systems before deployment	Reduces reliance on developer self-assessment.

Metric	Target	Reason
High-risk credential lifetime	Minutes or hours, not standing access	Limits persistence and blast radius.
Human nuclear authorisation	Lawful human decision plus independent human authentication/execution checks; no delegated AI authority	Prevents unilateral model or operator action.
Nuclear AI fail-safe review	Every two years and after material change or incident	Keeps assurance current as models and integrations change.
Critical-system manual fallback drill	At least annually	Tests continuity without AI services or AI-derived data.
Recovery exercise	At least annually for frontier labs and critical deployments	Tests whether control can actually be restored.

12 Conclusion

The policy objective is to retain effective human control over consequential systems as capability advances. The central unit of regulation is the deployed combination of model, tools, permissions and institutional authority. Narrow capability-based licensing, competent independent evaluation and enforceable deployment conditions can reduce exposure without licensing all ordinary AI use.

Nuclear safeguards should protect information integrity and deliberation as well as the final authorisation chain. Human presence alone is insufficient: operators need independent evidence, workable interfaces and rehearsed non-AI alternatives. Direct unauthorised launch remains a substantially more demanding pathway than peripheral compromise, but public evidence does not yield a reliable numerical bound on it.

International cooperation should begin with verifiable functions and protected evidence exchange, including China and other relevant states on reciprocal terms. National enforcement and a limited international safety institution are complementary. Neither a universal shutdown mechanism nor declarations without scrutiny can substitute for layered controls.

The proposed rules are a research and negotiation agenda. Their effectiveness should be tested through independent evaluations, incident learning and exercises, with thresholds and assumptions revised as evidence changes. The planning bands in section 8 are a prompt for robustness testing, not a warrant for confidence in a predicted date or probability of catastrophe.

Appendix A Proposed Treaty Principles

1. States retain responsibility for AI systems developed, deployed, or materially supported within their jurisdiction.

2. Frontier developers and designated compute providers must register and report high-risk activity to competent national authorities.
3. Training and deployment above agreed capability and risk thresholds require an approved safety case and independent evaluation.
4. Parties must protect frontier model weights and critical infrastructure according to common security tiers.
5. Serious AI incidents and near misses with cross-border or catastrophic potential must be reported through protected channels.
6. Retain lawful human nuclear-use authority and independent human checks; prohibit delegation to AI and automated retaliation without a fresh human decision.
7. Parties must maintain meaningful human control over nuclear and other strategic weapons and preserve non-AI fallback procedures.
8. Parties must conduct periodic AI-informed nuclear fail-safe reviews and provide agreed confidential attestations of completion.
9. Use protected declarations, process evidence and negotiated inspection rights, acknowledging incomplete verification and a defined anomaly-response process.
10. Non-compliance may result in corrective orders, suspension of designated compute or model services, export restrictions, financial penalties, and collective measures proportionate to the violation.
11. Safety information shared with the international agency must receive strong legal, technical, and diplomatic protection.
12. Thresholds and technical annexes must be reviewed at least annually so that regulation tracks capability rather than obsolete hardware measures.

Appendix B Minimum Frontier Safety Case

Section	Required evidence
Governance	Legal owner, accountable executive, risk committee, independent challenge, whistleblower route, regulator contact.
System boundary	Model, versions, tools, memory, data, credentials, infrastructure, users, jurisdictions, and external dependencies.
Threat model	Misuse, misalignment, insiders, model theft, supply chain, tool compromise, monitor compromise, and correlated failure.
Capability	Standardised and bespoke evaluations, uncertainty, elicitation method, external replication, and scaling forecast.

Section	Required evidence
Propensity	Deception, reward hacking, unauthorised action, oversight evasion, shutdown behaviour, and evaluation awareness.
Controls	Isolation, identity, authorisation, logging, monitoring, rate and resource limits, data provenance, rollback, and shutdown.
Residual risk	Probability and severity ranges, assumptions, dissenting expert views, red-team findings, and unresolved anomalies.
Operations	Deployment envelope, change control, monitoring thresholds, incident response, recovery, recall, and communication plan.
Assurance	Independent evaluator report, regulator findings, remediation evidence, and time-limited authorisation.

Appendix C Model Nuclear-AI Red Line Clauses

Article 1 Scope

AI system means a machine-based system that infers from inputs how to generate outputs such as predictions, content, recommendations or decisions that influence physical or virtual environments; systems vary in autonomy and adaptiveness. This functional definition draws on Article 3 of the EU AI Act.[5] A nuclear-relevant AI system materially influences nuclear warning, intelligence, readiness, communications, targeting, logistics or decision support. The authority prohibition applies to the authorisation and execution chain; ancillary applications remain subject to proportional safety controls. Article 2 also covers conventional automation to prevent a legacy-system loophole.

Compliance evidence: a maintained system inventory and documented classification decisions, accessible to the designated cleared authority.

Article 2 Prohibition on Autonomous Nuclear Authority

No State Party shall delegate nuclear-use authority to AI, whether exclusively or jointly, or permit AI to issue, alter or authenticate a nuclear-use order or replace a required human authorisation. Nuclear use shall require a fresh decision by the lawful human authority and positive, independently authenticated human checks in authorisation and execution. Automated initiation or retaliation without that decision is prohibited, including rules-based systems. Independently validated safety interlocks may inhibit an unsafe action but may not initiate nuclear use.

Compliance evidence: protected authorisation architecture, independently witnessed human-authentication tests and evidence that conventional automated initiation is excluded.

Article 3 Human Control and Fallback

Each State Party shall maintain tested non-AI warning assessment, communications, authentication and decision-support pathways. AI-derived information shall be identifiable and shall not be independently sufficient evidence for nuclear use. Operators shall be able to isolate AI components through independent controls while preserving essential safety functions. Interfaces and procedures shall support corroboration, dissent and task-validated deliberation. Automated escalation windows shall not bypass human verification; implementation shall preserve safe inhibit options and shall not impose untested delays.

Compliance evidence: witnessed isolation and fallback exercises, human-factors assessments and records of handling conflicting information.

Article 4 Fail-Safe Review

Each State Party shall complete a cleared, independent AI-informed fail-safe review at intervals not exceeding twenty-four months and after material changes or serious incidents. Reviews shall include indirect information and supply-chain dependencies. Annual non-AI drills shall test predetermined restoration, continuity, authentication and decision-quality criteria. Material failures shall receive named remediation owners, deadlines, proportionate interim operating restrictions and independent retesting.

Compliance evidence: review dates, independent reviewer credentials, drill results and verified closure of material findings.

Article 5 Negative Assurances and Verification

Each State Party shall attest annually to compliance with Articles 2–4 and provide agreed protected evidence. Routine international inspection shall concern declared civilian facilities within the instrument's authority. Classified assurance shall combine independent national review, sealed reports and managed demonstrations. Material anomalies, including inconsistent attestations, unexplained permission changes or missing required logs, shall trigger confidential consultation and a time-bounded response. Challenge access requires agreed safeguards. Attestation alone shall not be described as comprehensive verification.

Compliance evidence: signed declarations, evaluator scope and limitations, custody records and documented resolution of anomalies; protect operational secrets.

Article 6 Incidents and Crisis Communication

Parties shall maintain authenticated crisis contacts. A credible imminent risk of false nuclear warning, unauthorised action or cross-border escalation requires immediate notice through that channel when practicable. Submit an initial incident notification within twenty-four hours of recognition, a preliminary assessment within seventy-two hours, and continuing updates; target a fuller review within thirty days or explain the outstanding investigation. Early reports may be incomplete and shall distinguish facts, hypotheses and attribution confidence. These are proposed treaty deadlines, not a statement of universal existing law.

Compliance evidence: exercised crisis contacts, timestamped notifications and corrections, and documented investigations distinguishing uncertainty from attribution.

Article 7 Confidentiality and Compliance

Information exchanged under these clauses shall receive agreed legal and technical protection and shall not be used to obtain commercial advantage or unrelated military intelligence. Compliance concerns shall first be referred to a joint technical commission. Persistent non-compliance may lead to enhanced consultation, managed-access requests, suspension of cooperation, or proportionate national measures consistent with international law.

Compliance evidence: access controls, retention and deletion rules, an independent complaints route and records of proportionate, reviewable compliance decisions.

Appendix D Sources and Notes

- [1] Bengio, Y. et al., International AI Safety Report 2026, DSIT 2026/001, 3 February 2026, <https://internationalaisafetyreport.org/publication/international-ai-safety-report-2026>
- [2] OpenAI, The Hugging Face incident and other third-party impact from misaligned models, updated September 2026, <https://openai.com/hugging-face-incident-and-misalignment/>
- [3] OpenAI, Pacing model development in an era of cyber-critical capabilities, 18 August 2026, <https://openai.com/index/pacing-model-development-cyber-capabilities/>
- [4] National Institute of Standards and Technology, Artificial Intelligence Risk Management Framework 1.0, January 2023, and Generative Artificial Intelligence Profile NIST AI 600-1, July 2024, <https://www.nist.gov/itl/ai-risk-management-framework>
- [5] Regulation EU 2024/1689 laying down harmonised rules on artificial intelligence, Official Journal of the European Union, <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>
- [6] Council of Europe, Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, opened for signature 5 September 2024, <https://www.coe.int/en/web/artificial-intelligence/the-framework-convention-on-artificial-intelligence>
- [7] UK Department for Science, Innovation and Technology and Republic of Korea, Frontier AI Safety Commitments, AI Seoul Summit 2024, updated 7 February 2025, <https://www.gov.uk/government/publications/frontier-ai-safety-commitments-ai-seoul-summit-2024>
- [8] Williams, I., AI Is Moving Into Nuclear Weapons Systems. Fail-Safe Reviews Are Critical, Nuclear Threat Initiative, 17 September 2026, <https://www.nti.org/risky-business/ai-is-moving-into-nuclear-weapons-systems-fail-safe-reviews-are-critical/>

- [9] Hruby, J. and Miller, M. N., Assessing and Managing the Benefits and Risks of Artificial Intelligence in Nuclear-Weapon Systems, Nuclear Threat Initiative, 26 August 2021, <https://www.nti.org/analysis/articles/assessing-and-managing-the-benefits-and-risks-of-artificial-intelligence-in-nuclear-weapon-systems/>
- [10] International Atomic Energy Agency, Computer Security for Nuclear Security, IAEA Nuclear Security Series No. 42-G, Vienna, 2021.
<https://www.iaea.org/publications/13629/computer-security-for-nuclear-security>
- [11] International Atomic Energy Agency, Computer Security Techniques for Nuclear Facilities, IAEA Nuclear Security Series No. 17-T Rev. 1, Vienna, 2021.
<https://www.iaea.org/publications/14729/computer-security-techniques-for-nuclear-facilities>
- [12] National Institute of Standards and Technology, Cybersecurity Framework 2.0, February 2024,
<https://www.nist.gov/cyberframework>
- [13] Cyberspace Administration of China and six other authorities, Interim Measures for the Management of Generative Artificial Intelligence Services, effective 15 August 2023,
https://www.cac.gov.cn/2023-07/13/c_1690898327029107.htm
- [14] Cyberspace Administration of China and three other authorities, Measures for Labelling Artificial Intelligence Generated and Synthetic Content, effective 1 September 2025,
https://www.cac.gov.cn/2025-03/14/c_1743654684782215.htm
- [15] Cyberspace Administration of China and other authorities, Interim Measures for the Management of Anthropomorphic Artificial Intelligence Interactive Services, 10 April 2026,
https://www.cac.gov.cn/2026-04/10/c_1777558395078289.htm
- [16] Reuters, How China is preparing for the risk of AI escaping human control, 14 September 2026. Reporting on policy developments; author bylines and later revision dates are not relied upon.
<https://www.reuters.com/legal/litigation/how-china-is-preparing-risk-ai-escaping-human-control-2026-09-14/>
- [17] White House, America's AI Action Plan, July 2025,
<https://www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf>
- [18] United Nations General Assembly, Resolution 79/325: Terms of reference and modalities for the establishment and functioning of the Independent International Scientific Panel on Artificial Intelligence and the Global Dialogue on Artificial Intelligence Governance, 2025.
<https://docs.un.org/en/A/RES/79/325>
- [19] National Institute of Standards and Technology, Center for AI Standards and Innovation, model evaluations and standards work, accessed 19 September 2026,
<https://www.nist.gov/caisi>
- [20] Stanford Institute for Human-Centered Artificial Intelligence, AI Index Report 2026,
<https://hai.stanford.edu/ai-index/2026-ai-index-report>
- [21] International Atomic Energy Agency, Safeguards and Verification,
<https://www.iaea.org/topics/safeguards-and-verification>
- [22] Financial Action Task Force, The FATF Recommendations,
<https://www.fatf-gafi.org/en/publications/Fatfrecommendations/Fatf-recommendations.html>

- [23] OpenAI, The Hugging Face incident and the road ahead, 26 August 2026. Developer account of the July incident.
<https://openai.com/index/hugging-face-incident-and-the-road-ahead/>
- [24] METR and Redwood Research, Investigation of the OpenAI Hugging Face incident, 26 August 2026. External investigation with a bounded evidence-access period.
<https://metr.org/blog/2026-08-26-openai-hugging-face-incident-investigation/>
- [25] Anthropic, Investigating three real-world incidents in our cybersecurity evaluations, 30 July 2026. Developer disclosure.
<https://www.anthropic.com/news/investigating-incidents-cybersecurity-evals>
- [26] California, SB 53, Transparency in Frontier Artificial Intelligence Act, Chapter 138, approved 29 September 2025. Enacted statutory text.
https://leginfo.legislature.ca.gov/faces/billTextClient.xhtml?bill_id=202520260SB53
- [27] European Union, Regulation (EU) 2026/1744 (AI Omnibus); European Commission, AI Omnibus enters into force, 27 July 2026. Revised high-risk implementation timetable and oversight provisions.
https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=OJ:L_202601744
- [28] White House, National Security Presidential Memorandum/NSPM-11, 5 June 2026. Executive national-security policy.
<https://www.whitehouse.gov/presidential-actions/2026/06/national-security-presidential-memorandum-nspm-11/>
- [29] White House, Readout of President Joe Biden’s Meeting with President Xi Jinping of the People’s Republic of China, 16 November 2024.
<https://bidenwhitehouse.archives.gov/briefing-room/statements-releases/2024/11/16/readout-of-president-joe-bidens-meeting-with-president-xi-jinping-of-the-peoples-republic-of-china-3/>
- [30] Nuclear Threat Initiative, Advancing Nuclear Fail-Safe, project overview describing the US review completed in 2024.
<https://www.nti.org/about/programs-projects/project/advancing-nuclear-fail-safe/>
- [31] METR, Time Horizons, methodology and results; see also Kwa et al., Measuring AI Ability to Complete Long Tasks, arXiv:2503.14499 (2025).
<https://metr.org/time-horizons/>
- [32] IETF, RFC 9334, Remote ATtestation procedureS (RATS) Architecture, January 2023.
<https://www.rfc-editor.org/rfc/rfc9334.html>
- [33] National Telecommunications and Information Administration (2024). Dual-Use Foundation Models with Widely Available Model Weights Report. 30 July. Historical policy assessment, not a statement of current universal law.
<https://www.ntia.gov/programs-and-initiatives/artificial-intelligence/open-model-weights-report>
- [34] Chandra, B., Dunietz, J., Roberts, K., Lee, Y., Fontana, P. and Awad, G. (2024). Reducing Risks Posed by Synthetic Content: An Overview of Technical Approaches to Digital Content Transparency. NIST AI 100-4.
<https://doi.org/10.6028/NIST.AI.100-4>

- [35] Goddard, K., Roudsari, A. and Wyatt, J. C. (2012). Automation bias: a systematic review of frequency, effect mediators, and mitigators. *Journal of the American Medical Informatics Association*, 19(1), 121–127. Published online 16 June 2011.
<https://doi.org/10.1136/amiajnl-2011-000089>
- [36] Du, Y., Li, S., Torralba, A., Tenenbaum, J. B. and Mordatch, I. (2023). Improving Factuality and Reasoning in Language Models through Multiagent Debate. arXiv:2305.14325. Experimental research; not a nuclear-safety validation.
<https://arxiv.org/abs/2305.14325>
- [37] Mosier, K. L. and Skitka, L. J. (1999). Automation Use and Automation Bias. *Proceedings of the Human Factors and Ergonomics Society Annual Meeting*, 43(3), 344–348.
<https://doi.org/10.1177/154193129904300346>

Evidence Base and Limitations

Evidence cut-off: 20 September 2026. Official legal texts establish obligations only within their jurisdiction and operative dates. Developer disclosures are primary accounts by interested parties; external investigations improve scrutiny but remain limited by access and selection. Technical evaluations establish performance within tested conditions. Policy commentary supplies arguments, not incident-frequency estimates. High confidence denotes strong convergent support; moderate confidence credible but incomplete support; low confidence fragmented or assumption-sensitive support; very low confidence severe limitations. These labels are separate from event likelihood. No numerical band in section 8 has been independently calibrated. This is an authored policy research paper, not an experimentally validated prediction or a formally peer-reviewed journal publication.

Appendix E Implementation and Assurance Protocol

This annex is proposed policy language for adapting the paper to national regulation, a national-security directive or a negotiated treaty instrument. It creates no present legal obligation. Its requirements apply to covered Tier 3 activities and permissible Tier 4B applications; prohibited Tier 4A nuclear-use authority cannot be licensed through a safety case. The standards are origin-neutral and apply to open and closed systems according to capability, access and consequence.

E.1 Responsibilities and material changes

Original developer. Before a covered release, document model identity, evaluation scope, known limitations, plausible downstream enhancements and the irreversibility of distribution. Provide regulators and legitimate downstream recipients with proportionate safety information and an incident contact. Protect exploitable detail through controlled disclosure. A release decision should state which uses were evaluated and which remain outside the evidence.

Substantial modifier or fine-tuner. Record the parent model where known, version identifiers, relevant changes and the resulting evaluation evidence. A material change includes a demonstrated dangerous-capability increase, removal of an essential control, expansion of autonomous operation or adaptation to a higher-consequence use. Submit an amended or new

safety case before a covered release or deployment. Low-impact changes inside the approved envelope require documented regression checks; they do not automatically trigger full approval.

Integrator and operator. Identify the accountable legal entity for the deployed system. Evaluate the combination of model, retrieval sources, tools, memory, credentials, monitoring and human procedures. A model's previous approval is not authorisation for a new permission set or nuclear-relevant application. Maintain a change register, named incident owner, tested isolation procedure and a current operating envelope.

Distributor and infrastructure provider. Where legislation covers their activity, retain the minimum records needed to identify regulated operators and authorisations, preserve relevant evidence and respond to lawful, reviewable orders. Do not impose universal surveillance of ordinary users or assume that every passive repository can evaluate every uploaded model. Duties must be proportionate to actual control, scale, knowledge and risk.

Accountability and evidence. The national regulator should inspect release assessments, lineage declarations, change logs, evaluation reports and operating permissions. Allocate liability by the duty breached, causation, foreseeability and the actor's control, with applicable procedural protections. Contracts may allocate work but must not erase statutory duties. A missing or misleading lineage declaration is an assurance deficiency; a model hash is an identifier, not proof of safety or complete ancestry.

E.2 Cross-border enforcement and IFASA

A model hosted abroad or downloaded from an overseas laboratory remains subject to the importing or deploying state's applicable rules when used in a covered domestic activity. Participating states should designate a competent authority, recognise evaluations only within their tested scope, and exchange protected notices concerning serious failures. IFASA should support common methods, evaluator accreditation, incident coordination and technical consultation. Enforcement powers must arise from domestic law and the agreed treaty mandate.

For a covered operator, a substantiated deficiency should permit the national authority to require reassessment, narrow permissions, suspend the affected deployment or impose proportionate penalties. Emergency restrictions require documented grounds and prompt review. Findings about a foreign developer should distinguish verified evidence from allegations and provide a correction or challenge process. Nationality alone is not evidence of non-compliance.

For copied weights outside participating jurisdictions, available measures are necessarily incomplete: voluntary cooperation, restrictions on regulated domestic uses and services, protected incident exchange, and lawfully authorised measures against identifiable harmful conduct. No notice, licence condition or registry entry guarantees deletion of distributed weights. Report the residual enforcement gap rather than claiming worldwide recall or an international shutdown power.

E.3 Crisis information integrity

Accountable actor and trigger. The responsible nuclear authority shall apply a documented information-integrity procedure whenever synthetic, disputed or AI-summarised information could materially affect crisis assessment or nuclear posture. A designated duty officer shall retain responsibility for verification; an AI classifier shall not be the final arbiter of authenticity.

Required controls. Identify source, time, processing history where available, authentication status and unresolved contradictions. Separate observed facts from allegations and inference. Use independently sourced corroboration and authenticated official channels. Repetition across outlets or agreement among models using the same input shall not count as independent confirmation. Preserve original material and avoid summaries that remove uncertainty. Authentic origin does not guarantee factual accuracy; absent provenance does not prove falsification.

Assurance evidence. In protected exercises, test fabricated and genuine material, impersonation, unavailable provenance, compromised trusted sources and degraded communications. Record false acceptance, false rejection, verification latency, missed contradictions and resulting decision errors. Predetermine acceptable performance for each function and crisis condition through the sector safety case; do not invent one universal delay or accuracy threshold. Test against an independently established exercise ground truth, not a model-generated answer key.

Failure and response. Failure of an essential criterion requires proportionate restrictions on the affected AI-supported information pathway, restoration of validated alternatives, remediation and witnessed retesting. Preserve warning and safe inhibit functions. No unauthenticated video, synthetic-media score or AI-generated synthesis shall alone satisfy a prerequisite for nuclear use. This safeguard does not justify ignoring other independently corroborated evidence.

E.4 Meaningful human control and independent challenge

Required function. High-consequence AI-supported assessments shall permit an independently reasoned challenge, access to underlying evidence and a clear statement of unresolved disagreement. Qualified humans retain authority and responsibility. Where time and safety permit, record an initial human assessment before revealing AI recommendations. Use task-validated procedures when that sequence is impracticable; do not impose arbitrary delays that defeat protective action.

AI debate as an optional mechanism. A second model may support challenge only within an approved safety case. Document shared training dependencies where known, shared data, retrieval, tools and monitoring. Obtain initial assessments independently before any exchange, preserve dissent and test whether debate corrects errors or spreads them. Model agreement is not independent evidence, and an AI vote shall not determine nuclear-use authority. Different providers do not automatically establish independent failure modes.

Validation. Compare human-only, single-adviser and proposed multi-adviser arrangements in representative, protected exercises. Include misleading inputs, jointly wrong models, persuasive but unsupported explanations, contradictory evidence, AI unavailability and time pressure. Pre-register success criteria and record uncertainty. Measure false acceptance, missed hazards, decision quality, verification time, workload and recovery after withdrawal of AI support. Evaluate the complete operator-interface system, not benchmark accuracy alone.

Deployment gate and review. The cleared reviewing authority should permit an AI challenge configuration only where evidence supports a safety benefit without unacceptable delay or workload. Monitor for drift and repeat validation after material changes, at recurring fail-safe reviews and following serious incidents. Suspend or restrict a configuration that fails an essential criterion. Maintain a non-AI challenge route. These are proposed assurance

requirements; the cited studies do not establish a universal nuclear performance threshold.[35–37]

E.5 Compute notification and authorisation

Covered actors and triggers. Developers and providers of designated large-scale training services should report covered compute activity under published measurement rules. Capability and deployment triggers remain independently applicable, including to inexpensive fine-tuning and imported models. A spending figure such as USD 100 million shall not constitute a safe harbour for activity below that expenditure or conclusive evidence of dangerous capability above it.

Declarations and aggregation. Record the accountable entity, coordinated programme, participating providers, relevant compute estimates and uncertainty, material post-training activity, authorisation status and evaluation milestones. Define how coordinated runs and outsourcing are aggregated across locations and entities. Independently audit a risk-based sample and investigate discrepancies. Physical hardware ownership and commercial billing structure must not determine whether equivalent regulated activity escapes reporting.

Enforcement and review. Material concealment or evasion should trigger investigation and proportionate measures under law. Protect confidential information and allow correction of good-faith estimation errors. Review thresholds at least annually using observed capabilities, algorithmic efficiency and enforcement experience, with reasons and transition provisions for changes. Compute records support assurance; they cannot establish that every undeclared model or clandestine operation has been detected.

E.6 Adoption through a directive or treaty

A national-security directive should identify the issuing authority and legal basis, covered agencies and contractors, accountable executive, funding, reporting channel and cleared independent review body. An indicative adoption sequence is: designate responsibilities and inventory covered systems within ninety days; submit gap assessments and interim safeguards within one hundred and eighty days; complete initial representative exercises and remediation plans within twelve months. These are proposed planning milestones, subject to the jurisdiction's lawful authority and operational requirements.

A treaty annex should instead identify State Party duties, scope of consent, the relationship between IFASA and national authorities, protected evidence formats, review intervals, consultation rights and dispute procedures. It should specify how domestic implementation reaches relevant private actors, and how managed access is authorised. No clause should imply unlimited inspection of classified systems or legal authority over non-parties. Ratification, entry into force and enforcement arrangements require separate negotiation.

Review record. Each participating authority should maintain a protected register identifying the responsible actor, triggering activity, evidence submitted, decision, residual limitations, remediation deadline and next review. Publish aggregated implementation findings where safe. Independent reviewers should test whether controls change real operating practices rather than merely generating declarations. This annex supports auditable implementation; it does not convert incomplete evidence into proof of universal compliance.