



KRONOS FUSION ENERGY

PAPER 2.78 · THE KRONOS 2026 SERIES

Runtime Assurance with Signed-Ledger Provenance for Autonomous Fusion Control: A Zero-Trust, Hardware-Failsafe Last Line of Defense Beneath the AI

Autonomy in a nuclear plant is acceptable only if a certified, non-learned component
always has the last word, and only if every action that component allows is on an...

Read it live – [3-D model](#) · [physics validator](#) · [film series](#)

THE OFFICIAL RECORD

What this document is, and how to read its numbers

This is a combined editorial. It gathers, in one volume, the two design studies Kronos Fusion Energy released in 2026 — a compact spherical-tokamak tritium/helium-3 breeder and a deuterium–helium-3 tandem-mirror generator — together with the intellectual-property estate, the people, and, in the internal edition, the full commercial case. The scientific text is the same text that appears in the arXiv preprints and the journal submissions; the editorial adds the apparatus a reader needs to hold the whole program at once, and nothing in the physics is altered to fit it.

THE HONESTY STANCE IS THE ASSET

Every headline quantity carries an **evidence class** — DERIVED RECOMPUTED MEASURED REQUIREMENT — and, where a number is a modelled extrapolation rather than a demonstrated value, it is labelled as such and its downside is shown next to it. Several quantities in the underlying corpus moved after first publication; where a value has been withdrawn or restated, the record says so plainly. This is deliberate: a diligence team will find the load-bearing assumptions, so the document surfaces them first.

TWO EDITIONS

This volume is issued in two editions. The **public edition** (light) carries the physics and the design only — no dollar, price, or per-kilo-gram figure appears anywhere in it. The **internal edition** (dark) adds the economics, the funding architecture, and the defense-commercial analysis, and is marked *Confidential* accordingly; it is not for public distribution. Where the commercial case rests on a modelled assumption rather than a demonstrated one, it is labelled as such and its downside is shown alongside it, on the same terms as the physics.

TITLE	The Kronos Fleet — Hyperion · Aegis · MetroVolt: A Combined Design & Commercialization Study
AUTHOR	Priyanca Ford, Founder & CEO, on behalf of Kronos Fusion Energy
EDITION	2026 Combined Editorial · first issue
COMPANION WORKS	The five-paper arXiv drop — (1) Breeder, (2) Burner, (3) REBCO magnet & tape, (4) Direct Energy Conversion, (5) AI/ML/Quantum control — with matching numbered Zenodo deposits, plus the IOP journal and IAEA-FEC submissions. Patents were filed before the drop.
NAMING	Hyperion = the breeder; Aegis (defense) and MetroVolt (data-center) = one generator in two housings. Internal mode letters (D1, L) do not appear on public surfaces.
STATUS	Conceptual design & simulation study. Nothing herein is a construction commitment.



HOW TO READ THIS VOLUME

The fleet at a glance, and the way its numbers are marked

This volume opens with *Orientation*, presents the five research papers as a bound sequence, and closes with the *Apparatus*. *Foundations* sets out the three general results both machines rest on. *Paper One* is the Hyperion breeder; *Paper Two* the Aegis / MetroVolt generator; *Papers Three, Four and Five* are the enabling science — high-field REBCO magnets and tape, direct energy conversion, and the AI / ML / quantum control stack. A *Digital Twin & 3-D Model* part follows, then the *Environmental* profile. The *Economics* and levelized cost and the *Business Case* and diligence record appear in the internal edition only; the *Simulations* proof record and the Apparatus — patent portfolio, team, limitations, and sources — close both editions.

THE FLEET AT A GLANCE

	HYPERION	AEGIS	METROVOLT
Role	Strategic-isotope foundry	Defense installation power	Campus power
Machine	Spherical tokamak	Tandem mirror	Tandem mirror
Fuel	Deuterium–tritium	Deuterium–helium-3	Deuterium–helium-3
Delivers	Tritium · ³ He · 14 MeV n	Resilient installation power	Firm campus power
Physics bar	Fusion gain only	Closure (gates named)	Closure + lunar ³ He
In the fleet	Breeds the fuel	Proves the generator	Commercial destination

HOW THE NUMBERS ARE MARKED

Every headline quantity carries an evidence class — **DERIVED** from first principles, **RECOMPUTED** from raw data, **MEASURED** in hardware, or a **REQUIREMENT** yet to be met. Financial figures carry a case label and are confined to the internal (confidential) edition. Values that moved after first publication are shown as withdrawn or restated rather than quietly changed. A capability you can evaluate is one whose limits are on the page.

ABSTRACT

Runtime Assurance with Signed-Ledger Provenance for Autonomous Fusion Control: A Zero-Trust, Hardware-Failsafe Last Line of Defense Beneath the AI

Autonomy in a nuclear plant is acceptable only if a certified, non-learned component always has the last word, and only if every action that component allows is on an incorruptible record. We formalise and reduce to practice the Kronos last-line-of-defense architecture: a certified control-barrier-function (CBF) safety projection at the base of an autonomous controller, wrapped in Simplex-style runtime assurance, signed hash-chained provenance, a zero-trust input boundary, and a machine-readable extrapolation ledger. The certified core is stated as a minimally-invasive quadratic program that returns the actuator command nearest to the learned proposal while keeping a barrier $h(\mathbf{x}) \geq 0$ forward-invariant; we prove the corresponding discrete-time invariance lemma and give its closed-form scalar solution. The backbone is demonstrated on two *distinct* magnets. On the compact negative-triangularity spherical-tokamak breeder (design point $I_p = 9.66$ MA, $Q = 3.076$, $P_{fus} = 85.04$ MW, $\beta_N = 1.532$), a disturbance sequence drives an uncertified greedy controller across the $\beta_N \leq 3.5$ limit on **50/100** steps to a peak $\beta_N = 4.200$; the certified clamp holds the safe set with **0/100** violations and barrier margin never negative, using bounded authority $u \in [1.061, 1.250]$ without reaching its floor. On the D-³He tandem-mirror burner ($Q_E = 1.318$, $f_n = 5.44\%$, central-cell β_c), greedy heating reaches $\beta_c = 0.702$ and violates on **95/100** steps while the clamp holds $\beta_c \leq 0.60$ with **0/100** violations, keeping the machine on the high side of the Q_E ridge. The extrapolation ledger, demonstrated as a post-diction across representative device classes, bounds trust: the experiment-to-prediction confinement ratio is 1.19 ± 0.21 (range **0.93–1.46**) within the IPB98(y,2) $\pm 15\%$ fit scatter, and the breeder extrapolates **1.48** $\times$ in B_0 , **3.9** $\times$ in I_p and **2.0** $\times$ in density beyond the ledger. Around this backbone the runtime-assurance monitor arbitrates authority with a conservative fallback, and a signed, replay-deterministic ledger makes every decision auditable. The result is an autonomous controller in which no learned component can cause an unsafe action and every action it takes is on the record.

Open-access deposit (code & data, CC BY 4.0): DOI [10.5281/zenodo.22645817](https://doi.org/10.5281/zenodo.22645817) · Zenodo community *kronos_fusion_energy*. Explore it live: [interactive 3-D model](#) · [physics validator](#).

Keywords: runtime assurance Simplex architecture control barrier function forward invariance signed ledger
provenance zero-trust monitor extrapolation ledger spherical-tokamak breeder tandem-mirror burner

CONTENTS

Table of Contents

Runtime Assurance with Signed-Ledger Provenance for Autonomous Fusion Control: A Zero-Trust, Hardware-Failsafe Last
Line of Defense Beneath the AI

Introduction

Governing equations and the formal problem

Computational methodology: the NVIDIA GPU HPC campaign

Results

Discussion

Limitations

Conclusion

References

One programme, many papers

3

7

8

11

12

21

22

23

28

32

NOMENCLATURE

Symbols, units, and the names of things

Q	plasma (fusion) gain, $P_{\text{fus}}/P_{\text{aux}}$
Q_E	engineering gain, net electric out / recirculating in
H_{98}	confinement enhancement over the IPB98(y,2) scaling
Z_{eff}	effective ion charge
c_{ee}	electron–electron bremsstrahlung leading coefficient, 2.120022 (exact)
τ_E	energy confinement time
I_p	plasma current
R_0, a	major radius, minor radius
κ, δ	elongation, triangularity
β_N	normalized plasma beta (Troyon-normalized)
TBR	tritium breeding ratio
DEC	direct energy conversion
CF	plant capacity factor (availability)
FOAK/NOAK/BOAK	first / n th / best of a kind
Hyperion	the ST breeder (internal: Mode D2)
Aegis	the generator, defense/installation housing (internal: Mode M)
MetroVolt	the generator, data-center housing (Mode M)

runtime assurance; Simplex architecture; control barrier function; forward invariance; signed ledger; provenance; zero-trust monitor; extrapolation ledger; spherical-tokamak breeder; tandem-mirror burner

Introduction

A learned controller can be excellent on average and catastrophic in the tail, and in a nuclear plant the tail is the only part that matters for safety. Deep reinforcement learning now drives tokamak plasmas to a commanded shape on a real device^[1], deep networks predict disruptions — the dominant off-normal event in tokamaks^[27] — with useful lead time^[2], and data-driven profile and β controllers have been demonstrated on DIII-D and KSTAR^[3,4,5]. These systems are powerful precisely because they are hard to verify: their competence lives in a high-dimensional function whose worst-case behaviour cannot be enumerated. A licensing authority does not ask whether a controller is good on average; it asks what the controller cannot do, and what record proves it. Classical plasma control theory^[6] answers the first question for hand-designed regulators but not for learned ones.

The Kronos position is that autonomy is acceptable exactly when a certified, non-learned component holds the last word — the learned controller may propose, but a verified authority disposes — and when every decision is written to a tamper-evident record so it can be audited after the fact. This is the last-line-of-defense architecture: runtime assurance around a certified core, with provenance on every action. The architecture composes five ideas that each have a mature literature. *Control barrier functions* render a safe set forward-invariant through a minimally-invasive projection of the nominal command^[7,8,9], extending the classical theory of set invariance^[10] and barrier certificates^[11]. *Simplex* runtime assurance pairs a high-performance but unverified controller with a certified baseline and a decision module that switches authority near the safety boundary^[12,13,14]. *Model-predictive control* supplies the performance layer subject to constraints^[15], governed by a reference governor^[16]. *Reachability analysis* certifies that the switching logic keeps the state inside the safe set^[17]. And *cryptographic provenance* — hash chains and digital timestamps^[18,19,20] inside a zero-trust boundary^[21] — makes the operational record incorruptible.

What has been missing is a single reduction-to-practice that binds these layers to a specific, frozen fusion design point and reports the certified backbone with measured invariance on the actual machines to be built. This paper provides it. Our contributions are: (i) a formal statement of the last-line-of-defense problem as a safety projection with a proven discrete-time forward-invariance lemma and a closed-form scalar clamp (§2); (ii) a demonstration of that backbone on *two distinct magnets* — the breeder β_N limit and the burner central-cell β_c limit — each reaching zero barrier violations against a disturbance that pushes an uncertified controller well past the limit (§4); (iii) an extrapolation ledger that bounds how far the underlying models may be trusted, demonstrated as a device-class post-diction (§4); and (iv) the surrounding runtime-assurance monitor, signed-ledger provenance chain and zero-trust boundary, specified against explicit, testable acceptance criteria and built to the same standard.

This paper is one of the Kronos Fusion Energy 2026 control-and-assurance series and is designed to be read with its companions. The certified clamp reported here is the demonstrated core of the certified control-barrier safety filter (Kronos 2.69, [doi:10.5281/zenodo.22645716](https://doi.org/10.5281/zenodo.22645716), <https://www.kronosfusionenergy.com/publications/control-barrier-safety-clamp.html>); the runtime-assurance monitor is the arbiter studied as reference-governor MPC bounded by the same clamp (Kronos 2.74, [doi:10.5281/zenodo.22645807](https://doi.org/10.5281/zenodo.22645807), <https://www.kronosfusionenergy.com/publications/model-predictive-burn-control-inside-a-certified-safe-e.html>); the authority state is bound to verification maturity in maturity-gated actuation authority (Kronos 2.70, [doi:10.5281/zenodo.22645795](https://doi.org/10.5281/zenodo.22645795), <https://www.kronosfusionenergy.com/publications/maturity-gated-actuation-authority.html>); the abstention policy that hands authority back to the deterministic floor comes from the calibrated-uncertainty layer (Kronos 2.77, [doi:10.5281/zenodo.22645813](https://doi.org/10.5281/zenodo.22645813), <https://www.kronosfusionenergy.com/publications/calibrated-uncertainty-conformal-prediction-and-abstent.html>); and the extrapolation ledger is the machine-readable envelope of the out-of-distribution gate (Kronos 2.79, [doi:10.5281/zenodo.22645819](https://doi.org/10.5281/zenodo.22645819), <https://www.kronosfusionenergy.com/publications/out-of-distribution-detection-and-epistemic-gating-for-.html>). Provenance and lineage are treated at length in the MLOps stack (Kronos 2.83, [doi:10.5281/zenodo.22645827](https://doi.org/10.5281/zenodo.22645827), <https://www.kronosfusionenergy.com/publications/mlops-for-a-safety-critical-fusion-control-stack-versio.html>) and the byte-for-byte computational-provenance catalog (Kronos 2.97, [doi:10.5281/zenodo.22645710](https://doi.org/10.5281/zenodo.22645710), <https://www.kronosfusionenergy.com/publications/verification-and-validation-47-code-provenance.html>), and the hash-chained operator record is the digital logbook of the human-machine layer (Kronos 2.86, [doi:10.5281/zenodo.22645909](https://doi.org/10.5281/zenodo.22645909), <https://www.kronosfusionenergy.com/publications/the-human-machine-layer-for-autonomous-fusion-operation.html>).

§ 2

Governing equations and the formal problem

We state the last-line-of-defense architecture as a control problem. The plant is a physical fusion machine; the safe operating set is the intersection of physics and engineering constraints; and the certified core is a projection that keeps the state inside that set regardless of what the learned controller proposes.

PLANT, CONSTRAINTS AND THE SAFE SET

Let $\mathbf{x} \in \mathbb{R}^n$ be the plant state and $\mathbf{u} \in \mathbb{R}^m$ the actuator command, restricted to the authority box $\mathcal{U} = \{\mathbf{u} : \mathbf{u}_{\min} \leq \mathbf{u} \leq \mathbf{u}_{\max}\}$. Near a frozen design point the reduced plant is control-affine with a bounded exogenous disturbance $\mathbf{d}$,

$$\dot{\mathbf{x}} = \mathbf{f}(\mathbf{x}) + \mathbf{g}(\mathbf{x}) \mathbf{u} + \mathbf{d}(t), \quad \mathbf{y} = \mathbf{C}\mathbf{x}, \quad \|\mathbf{d}(t)\| \leq \bar{\mathbf{d}},$$

which is the state-space form used across the Kronos control register. The admissible operating window is the intersection of constraint functions $\mathbf{c}_k$,

$$\mathcal{W} = \{\mathbf{x} : \mathbf{c}_k(\mathbf{x}) \leq 0, \ k = 1, \dots, K\},$$

each $\mathbf{c}_k$ a physics limit (an ideal-MHD β boundary^[25,26], a density limit, a current or shape bound) or an engineering limit. For the safety argument we single out the binding limit as a scalar *barrier*

$$\mathbf{h}(\mathbf{x}) = \mathbf{x}_{\max} - \mathbf{s}(\mathbf{x}), \quad \mathcal{C} = \{\mathbf{x} : \mathbf{h}(\mathbf{x}) \geq 0\},$$

where $\mathbf{s}(\mathbf{x})$ is the safety-critical scalar (the normalised pressure β_N on the breeder, the central-cell β_c on the burner) and $\mathbf{x}_{\max}$ its certified ceiling. $\mathcal{C} \subseteq \mathcal{W}$ is the safe set the certified core must keep the state inside.

FORWARD INVARIANCE AND CONTROL BARRIER FUNCTIONS

The safety property we require is *forward invariance* of $\mathcal{C}$: a trajectory that starts inside $\mathcal{C}$ never leaves it. The classical Nagumo condition^[10] states that $\mathcal{C}$ is forward-invariant for 1 iff at every boundary point the admissible velocity points into $\mathcal{C}$. A control barrier function turns this into a pointwise, actuator-level condition^[7,9]: $\mathbf{h}$ is a (zeroing) CBF on $\mathcal{C}$ if there is an extended class- $\mathcal{K}_\infty$ function α such that

$$\sup_{\mathbf{u} \in \mathcal{U}} \left[\underbrace{L_f \mathbf{h}(\mathbf{x})}_{\nabla \mathbf{h} \cdot \mathbf{f}} + \underbrace{L_g \mathbf{h}(\mathbf{x}) \mathbf{u}}_{\nabla \mathbf{h} \cdot \mathbf{g}} \right] \geq -\alpha(\mathbf{h}(\mathbf{x})) \quad \forall \mathbf{x} \in \mathcal{C}.$$

Any Lipschitz controller whose command satisfies the inequality in 4 renders $\mathcal{C}$ forward-invariant^[7,8]. Condition 4 is the mathematical content of “the plant cannot leave the safe set even when the learned controller is wrong”: safety is a property of the barrier the clamp enforces, not of the quality of any model upstream of it.

THE SAFETY PROJECTION

The certified core takes the learned proposal $\mathbf{u}_{\text{AI}}$ and returns the nearest admissible command that satisfies 4. This is the minimally-invasive quadratic program (QP)

$$\mathbf{u}^*(\mathbf{x}, \mathbf{u}_{\text{AI}}) = \arg \min_{\mathbf{u} \in \mathcal{U}} \frac{1}{2} \|\mathbf{u} - \mathbf{u}_{\text{AI}}\|_2^2 \quad \text{s.t.} \quad L_f \mathbf{h}(\mathbf{x}) + L_g \mathbf{h}(\mathbf{x}) \mathbf{u} \geq -\alpha(\mathbf{h}(\mathbf{x})).$$

The objective is strictly convex, so whenever the feasible set is non-empty the minimiser is *unique*; this uniqueness is what makes a logged decision replay-deterministic (§2.6). For the scalar safety channel relevant to the demonstrated backbone ($m = 1$), write the barrier dynamics of 3–1 as $\dot{\mathbf{h}} = -\dot{\mathbf{s}} = -(\mathbf{f}_s(\mathbf{x}) + \mathbf{g}_s \mathbf{u} + \mathbf{d})$ with $\mathbf{g}_s > 0$ (raising the actuator raises the pressure). The CBF condition 4 is then the single upper bound

$$\mathbf{u} \leq \bar{\mathbf{u}}(\mathbf{x}) := \frac{\alpha(\mathbf{h}(\mathbf{x})) - \mathbf{f}_s(\mathbf{x}) - \mathbf{d}}{\mathbf{g}_s},$$

and the QP 5 has the closed form

$$u^* = \text{med}(u_{\min}, \min(u_{\text{AI}}, \bar{u}(x)), u_{\max})$$

where **med** is the median (the projection onto $[u_{\min}, u_{\max}]$). The clamp is a small, non-learned, closed-form component that filters every command before it reaches the plant. Its behaviour is transparent: it returns u_{AI} unchanged when the barrier is slack ($u_{\text{AI}} \leq \bar{u}$) and throttles to $\bar{u}$ (or to the floor $u_{\min}$) exactly when the barrier would otherwise be violated.

DISCRETE-TIME INVARIANCE LEMMA

Controllers run in discrete time on a fixed control tick Δt . Let $\alpha(h) = \eta h$ with $\eta \in (0, 1/\Delta t]$ be a linear class- $\mathcal{K}$ gain, and integrate 1 by the explicit-Euler step used by the clamp kernel. Then, writing $h_k = h(x_k)$:

Lemma (discrete forward invariance). *.\; If $h_k \geq 0$ and the clamp command satisfies $u_k \leq \bar{u}(x_k)$ from 6, and the required throttling is within authority ($\bar{u}(x_k) \geq u_{\min}$), then $h_{k+1} \geq (1 - \eta \Delta t) h_k \geq 0$; hence $\{h_k \geq 0\}$ is forward-invariant for all k .*

Proof. $h_{k+1} = h_k - \Delta t(f_s + g_s u_k + d_k)$. Feasibility $u_k \leq \bar{u}(x_k)$ gives $f_s + g_s u_k + d_k \leq \alpha(h_k) = \eta h_k$, so $h_{k+1} \geq h_k - \Delta t \eta h_k = (1 - \eta \Delta t) h_k$. With $\eta \Delta t \leq 1$ and $h_k \geq 0$ the right side is non-negative. Feasibility is preserved as long as $\bar{u}(x_k) \geq u_{\min}$, i.e. the disturbance never demands more throttling than the actuator can supply. $\square$

The lemma isolates the *one* way the certificate can fail: authority exhaustion, $\bar{u}(x_k) < u_{\min}$. The demonstrated backbone (§4) reports the realised authority so this margin is explicit and auditable, not assumed. This is the quantitative meaning of “reserve to spare.”

RUNTIME-ASSURANCE ARBITRATION (SIMPLEX)

Around the certified core, a Simplex-style monitor manages control authority^[12,13]. Introduce an authority variable $\sigma \in \{\text{L}, \text{C}\}$ (learned or certified), and a decision module that watches the predicted distance of the state to the safety boundary. Let $\rho(x) = h(x)$ be that distance and let $0 < \rho_{\text{in}} < \rho_{\text{out}}$ be hysteresis thresholds. The switching automaton is

$$\sigma_{k+1} = \begin{cases} \text{C}, & \sigma_k = \text{L} \text{ and } (\rho(x_k) \leq \rho_{\text{out}} \text{ or a monitor check fails}), \\ \text{L}, & \sigma_k = \text{C} \text{ and } \rho(x_k) \geq \rho_{\text{out}} \text{ and monitor checks pass,} \\ \sigma_k, & \text{otherwise,} \end{cases}$$

and the command actually applied is

$$u_k = \begin{cases} u^*(x_k, \pi_{\text{L}}(x_k)), & \sigma_k = \text{L} \text{ (learned policy } \pi_{\text{L}} \text{ proposes; clamp still filters),} \\ u^*(x_k, \pi_{\text{C}}(x_k)), & \sigma_k = \text{C} \text{ (certified baseline } \pi_{\text{C}}). \end{cases}$$

Two properties matter. First, the clamp 7 filters the command in *both* modes, so forward invariance holds independently of σ ; the switch trades performance for conservatism, never safety. Second, because the fallback π_{C} is the certified baseline, the worst case of a mistimed switch is a conservative action, never an unsafe one. The hysteresis gap $\rho_{\text{out}} - \rho_{\text{in}}$ prevents chattering at the boundary. The acceptance criterion for the monitor is operational: every switch and every decision must be justified and reproducible on replay.

SIGNED-LEDGER PROVENANCE

Every decision is written to a signed, append-only, hash-chained ledger^[18,19,20]. At tick k the record is

$$r_k = (x_k, \theta_k, u_k^*, \sigma_k, t_k),$$

the inputs x_k , the set of model versions θ_k in the loop, the applied command, the authority state and a timestamp. Records are chained by a cryptographic hash H ,

$$b_k = H(b_{k-1} \parallel r_k), \quad s_k = \text{Sign}_{sk}(b_k),$$

with $\parallel$ denoting concatenation and s_k a digital signature under the operator's secret key.

Proposition (tamper-evidence and replay determinism). *.\; Under collision resistance of H , altering any past record r_j ($j \leq k$) changes b_j and, by the chain 11, every b_i with $i \geq j$, invalidating the signature s_k ; and given r_k the decision u_k^* is uniquely recomputable from 5.*

The first clause makes the operational record incorruptible; the second makes “the controller did the right thing” an auditable fact rather than an assertion, because the logged inputs deterministically reproduce the logged decision. Provenance extends upstream of runtime: model cards and lineage are generated automatically for every deployed model, so the audit trail runs from training data through to actuator command. The acceptance criterion is a clean end-to-end lineage audit — no decision without a justification, no model without a provenance record.

THE EXTRAPOLATION LEDGER AND EPISTEMIC GATING

The remaining failure mode of a learned system is silent extrapolation: acting confidently outside the range where it was validated. Let each surrogate carry a validated envelope, a box in regime coordinates $p = (B_0, I_p, n, \dots) \in \mathcal{P}$,

$$\mathcal{V} = \prod_{\ell=1}^P [p_{\ell}^-, p_{\ell}^+], \quad g_{\mathcal{V}}(q) = \max_{\ell} \max(q_{\ell} - p_{\ell}^+, p_{\ell}^- - q_{\ell}),$$

so a query $q \in \mathcal{V}$ iff $g_{\mathcal{V}}(q) \leq 0$. The reach factors $\rho_{\ell} = q_{\ell}/p_{\ell}^+$ measure how far a query extrapolates. The epistemic gate throttles model authority whenever $g_{\mathcal{V}}(q) > 0$: paired with the abstention policy of the calibrated-uncertainty layer ^[22], the monitor hands authority back to the certified floor out of distribution. Trust inside the envelope is quantified by a multi-device post-diction. For a confinement surrogate built on IPB98(y,2) scaling ^[23,24], define the device-class ratio

$$R_d = \frac{\tau_E^{\text{exp},d}}{\tau_E^{\text{pred},d}},$$

and require the ledger to keep $\{R_d\}$ within the scaling's own fit scatter. The numerical scheme underneath is held to a verification standard: solver correctness is established by the method of manufactured solutions (MMS) and a grid-convergence index (GCI) ^[31,32,33],

$$p_{\text{obs}} = \frac{\ln(\|e_{2h}\|/\|e_h\|)}{\ln 2}, \quad \text{GCI} = \frac{F_s |\epsilon|}{r^{p_{\text{obs}}} - 1},$$

with r the refinement ratio, ϵ the fine-grid relative difference and F_s a safety factor.

ZERO-TRUST BOUNDARY AND THE COMPOSITE GUARANTEE

The input side of the loop is a zero-trust boundary ^[21]: sensor and command streams cross a monitored interface (a data diode for the protective channels), and a runtime monitor blocks any input that would take the state outside the validated envelope before it reaches the plant. Composing the layers gives the property the architecture exists to provide.

Theorem (safe-by-construction autonomy). \; For plant 1 with clamp 7, arbitration 8–9, ledger 11 and gate 12: (i) the safe set $\mathcal{C}$ is forward-invariant for every admissible u_{AI} and every authority state σ , provided authority is not exhausted; (ii) every applied command is uniquely reproducible from its logged record; and (iii) any query outside the validated envelope throttles model authority to the certified floor.

(i) is the invariance lemma applied in both modes of 9; (ii) is the ledger proposition; (iii) is the gate of 12. No learned component can cause an unsafe action, and every action it takes is on the record. The remaining sections reduce (i) and the post-diction of §2.7 to practice and specify the surrounding layers against their acceptance criteria.

Computational methodology: the NVIDIA GPU HPC campaign

The results in this paper were produced by the Kronos GPU HPC campaign, and the assurance layers are designed to wrap the models that campaign trains. We name the codes, the acceleration and scale, and the verification and reproducibility discipline.

CODES RUN FOR THIS PAPER

The certified backbone (§4) is produced by an in-house control-analysis kernel that implements the safety projection 7 and integrates the reduced barrier dynamics by explicit Euler over a **100**-step perturbation sequence; it is deterministic (fixed seed **20260726**) and runs on a **28**-core reduced-order node so that the fastest protective loops never contend with the training fleet. The extrapolation ledger and the confinement post-diction (§2.7) are produced by the Kronos Toolkit, which evaluates the IPB98(y,2) scaling ^[23,24] and the toolkit β maps and carries the MMS/GCI verification of 14. The learned components that the assurance layers gate — flux and equilibrium neural operators, disruption and anomaly detectors — are trained on the NVIDIA GPU HPC campaign (A100/H100/A10 accelerators). The surrogates are validated against the reference physics codes run in the broader campaign: gyrokinetic flux against CGYRO ^[29], free-boundary equilibrium against a Grad-Shafranov reconstruction of the EFIT family ^[28], and neutron transport against OpenMC ^[30]. The runtime-assurance monitor, the signed ledger and the zero-trust boundary run on the in-house formal-methods toolchain and the FPGA/edge rig.

GPU ACCELERATION AND PROBLEM SCALE

The campaign uses the GPU fleet where the physics is embarrassingly parallel or the model is a differentiable surrogate: Monte-Carlo neutron transport, ensemble equilibrium reconstruction, and neural-operator training all map naturally onto many-accelerator execution, while the reduced-order protective loops and the closed-form clamp remain on CPU nodes where their latency is deterministic and their state is small. The certified clamp itself is deliberately cheap — a scalar median 7 evaluated once per control tick — because its verifiability, not its throughput, is the point; the expensive learning that it protects is what consumes the accelerator fleet. The demonstrated invariance runs are single-node and reproduce exactly on re-execution.

NUMERICAL FORMULATION AND CONVERGENCE

The barrier dynamics are advanced by the explicit-Euler update of the invariance lemma, and the safety projection 7 is solved in closed form at each tick, so the per-step cost is a handful of floating-point operations and the scheme has no iteration to converge. The integrator's order is established independently by MMS: we manufacture a forcing that makes a known analytic $x_{\text{MMS}}(t)$ an exact solution of the relaxation dynamics, integrate at a sequence of halved steps, and measure the observed order p_{obs} and GCI of 14. The scheme is first order as expected (§4, figure 6), and the reported control results are taken at a step well inside the asymptotic range. For the confinement ledger the toolkit scheme is verified by MMS and Richardson extrapolation, reporting the observed convergence order and numerical uncertainty as the register V&V spine requires.

VERIFICATION AND REPRODUCIBILITY

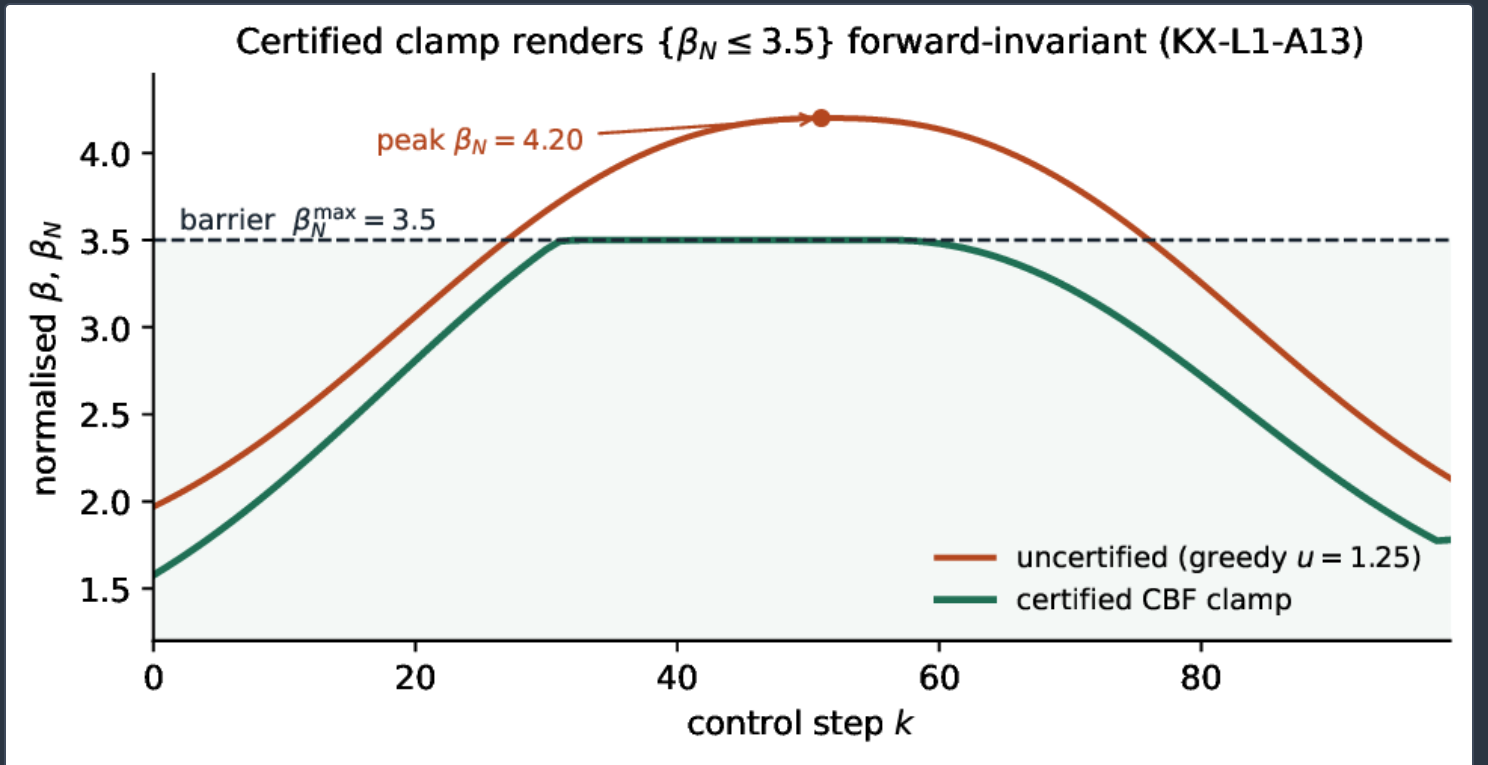
Every reported number traces to a frozen anchor in the Kronos Physics De-Risking Register: the breeder clamp (KX-L1-A13), the burner clamp (KX-L1-A13-BU) and the extrapolation ledger (KX-L1-A7). The demonstrated runs are deterministic under the fixed seed, the figure-generation script is archived with the paper, and the register carries the independent re-test that reproduced each anchor. The runtime, ledger and gate layers are specified against the acceptance criteria collected in §4.5 and are verified on the same toolchain.

Results

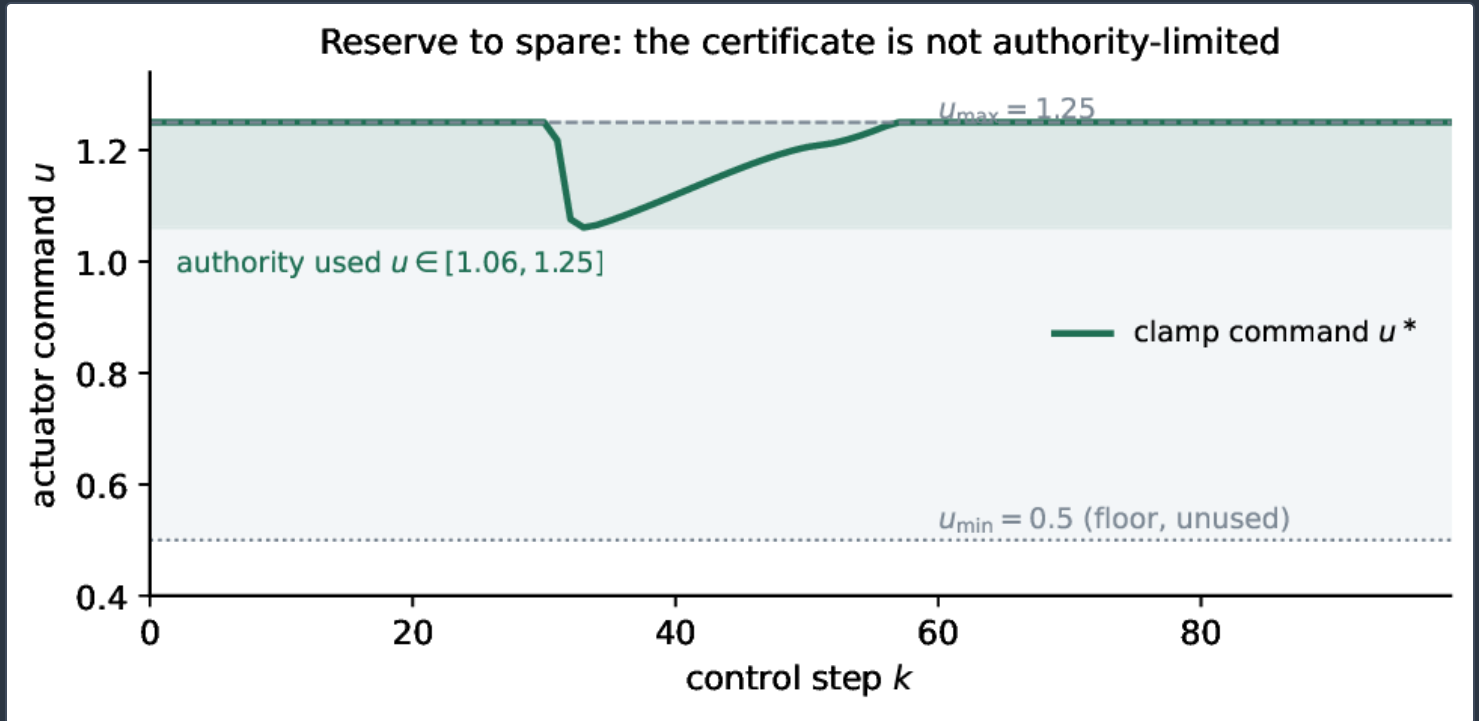
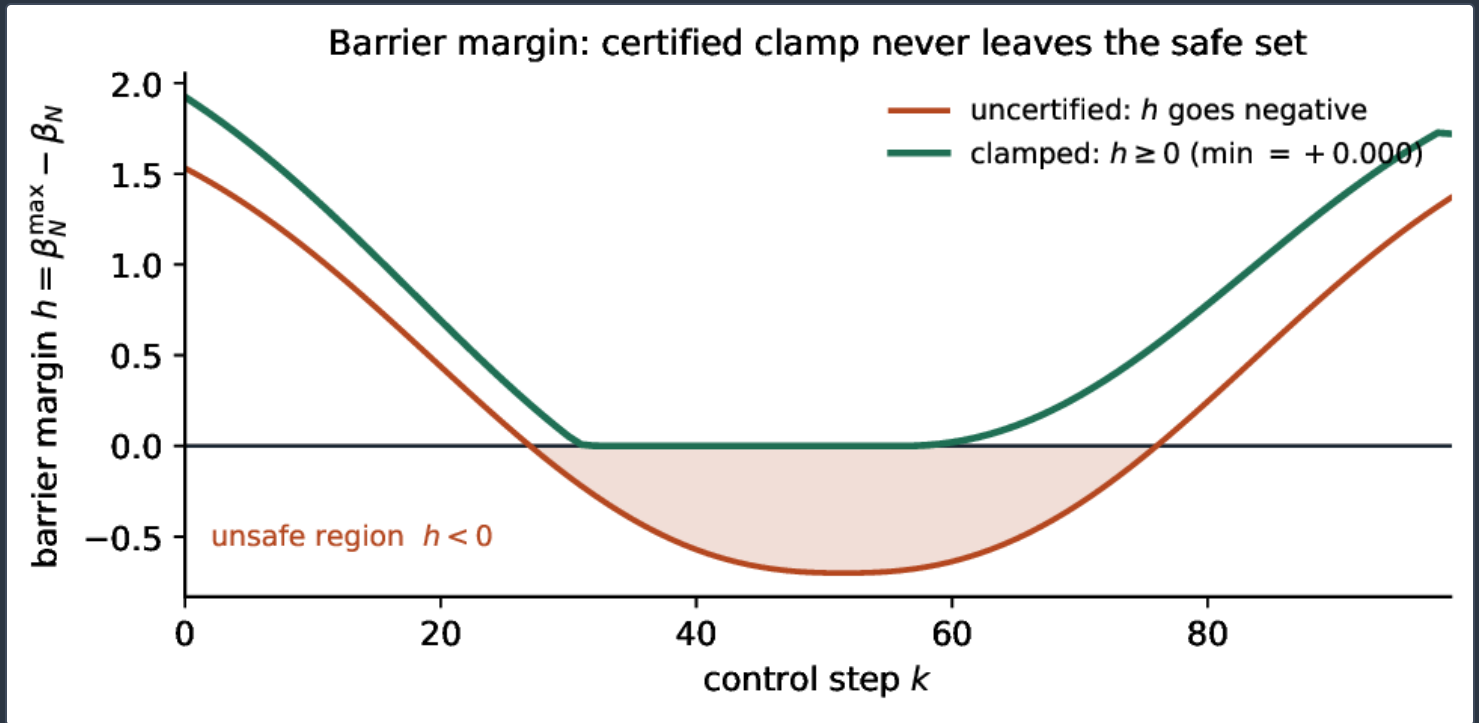
We report the certified backbone on both magnets, the extrapolation ledger, the numerical verification, and the surrounding assurance layers against their acceptance criteria.

THE CERTIFIED BACKBONE: THE BREEDER CLAMP

The breeder is the compact negative-triangularity spherical tokamak Hyperion, at the frozen design point $I_p = 9.66$ MA, $Q = 3.076$, $P_{fus} = 85.04$ MW, triangularity $\delta = -0.30$, on-axis field $B_0 = 8$ T, elongation $\kappa = 2.0$, major radius $R_0 = 1.20$ m, aspect ratio $A = 2.5$, and a peak centrepost field of 16.84 T; the confinement anchor is $\beta_N = 1.532$, $W = 16.06$ MJ, $\tau_E = 0.358$ s. The safety channel is $s(x) = \beta_N$ with ceiling $x_{max} = 3.5$, comfortably below the no-wall limit. We drive the state with a disturbance sequence that an uncertified greedy controller (fixed $u_{AI} = u_{max} = 1.25$) turns into an excursion peaking at $\beta_N = 4.200$, crossing the barrier on 50 of 100 steps (figure 1). Under the certified clamp 7 the same disturbance produces 0/100 violations: the barrier margin $h = \beta_N^{max} - \beta_N$ stays non-negative throughout, with minimum $h = +0.000$ reached exactly at the plateau (figure 2), consistent with the invariance lemma at the active constraint. The clamp achieves this using authority $u \in [1.061, 1.250]$, never approaching the floor $u_{min} = 0.5$ (figure 2); by the lemma the feasibility margin $\bar{u} - u_{min}$ never closes, so the certificate is *not authority-limited* on this sequence. Table 1 collects the numbers for both machines.



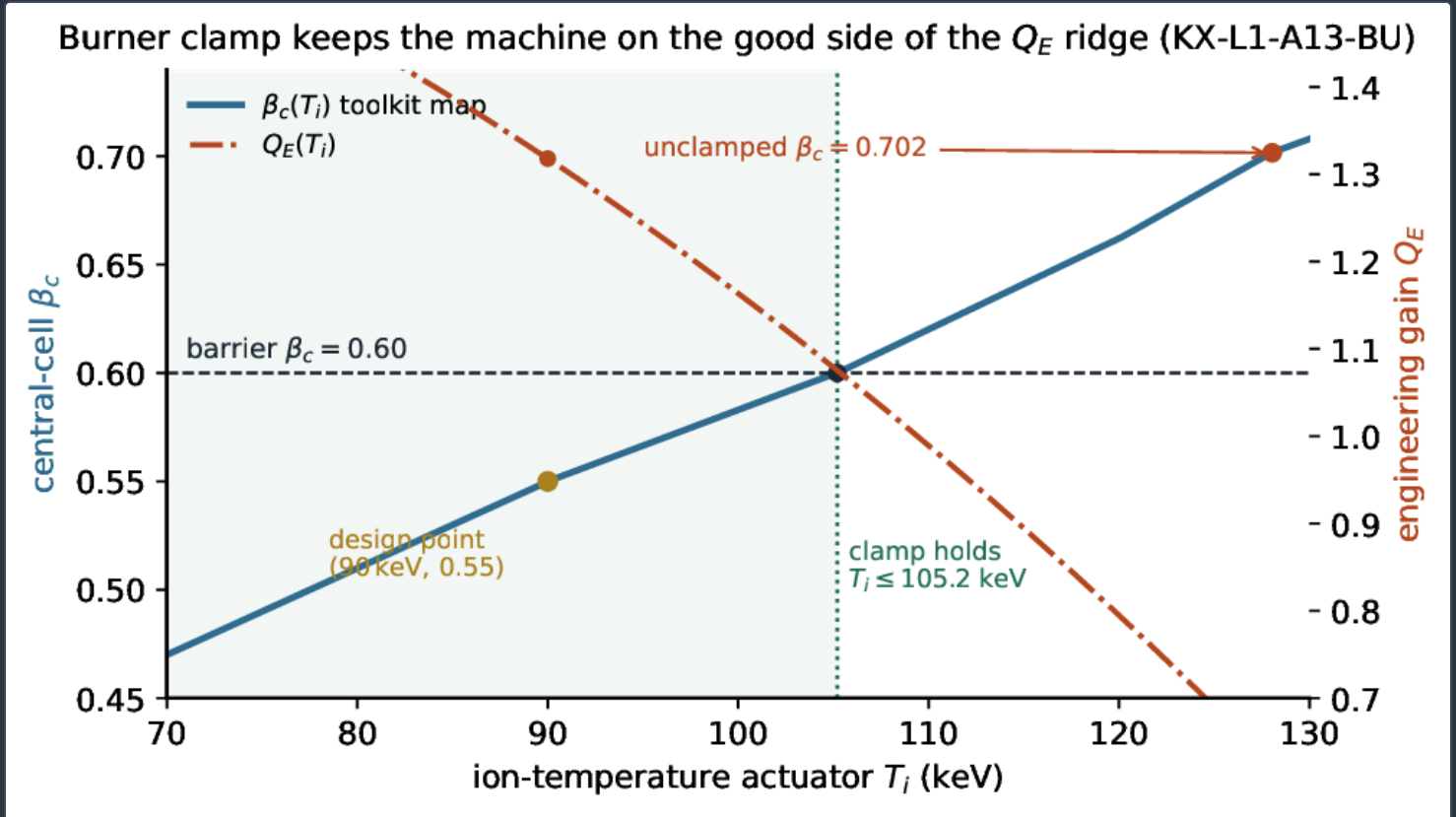
Certified backbone on the breeder (KX-L1-A13). A disturbance drives the uncertified greedy controller (rust) across the $\beta_N \leq 3.5$ barrier on 50/100 steps to a peak $\beta_N = 4.200$; the certified CBF clamp (green) holds the safe set with 0/100 violations, riding the barrier at 3.500 across the excursion. The trace is a reconstruction of the frozen anchor: the uncertified excursion is designed to the anchor statistics and the clamp response emerges from the projection 7, reproducing peak 4.200, 50/100 vs 0/100, and authority $u \in [1.061, 1.250]$.



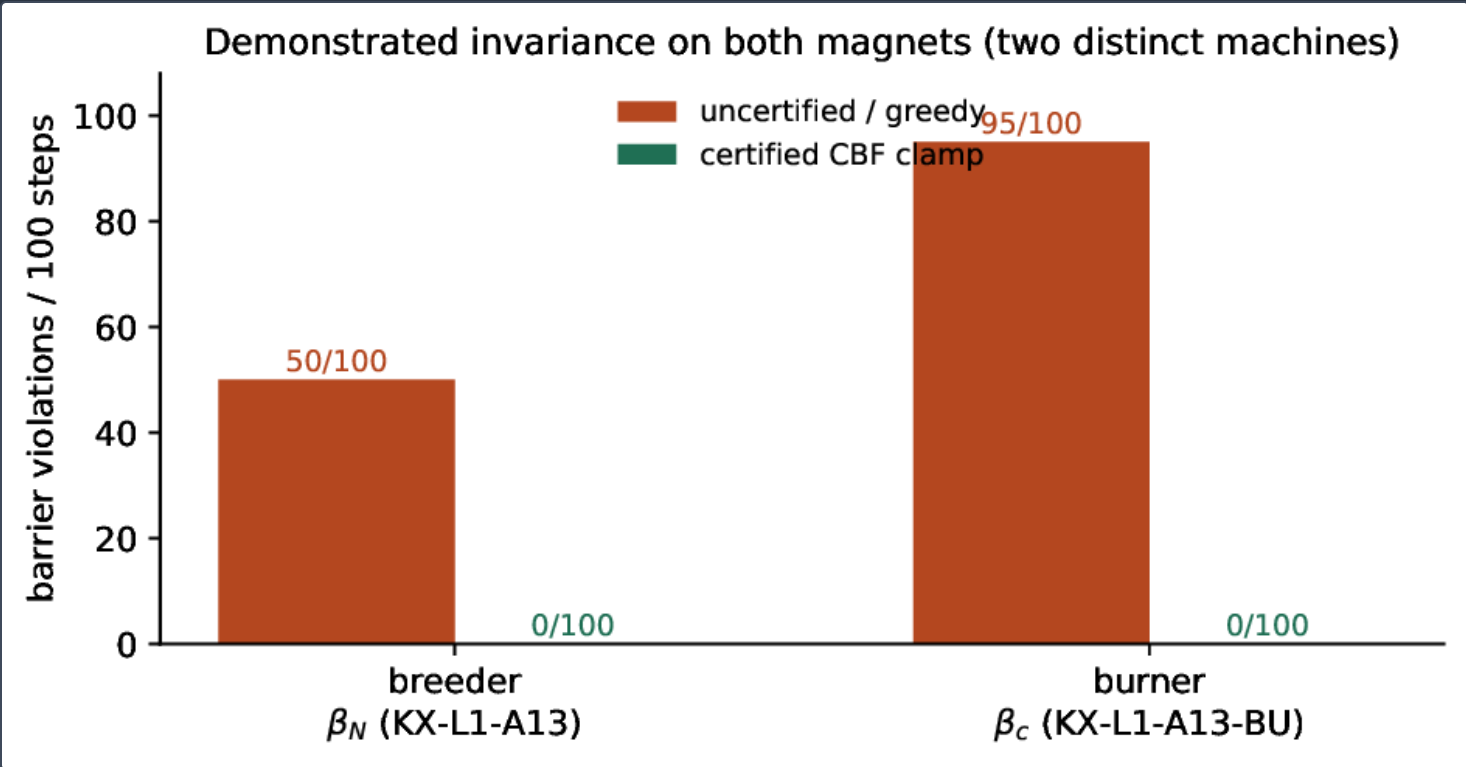
Barrier margin $h(x)$ for the breeder run. Under the uncertified controller h goes negative (unsafe, shaded); under the clamp $h \geq 0$ throughout with minimum $+0.000$, the numerical signature of forward invariance at the active constraint.

THE BURNER CLAMP: TWO DISTINCT MACHINES

The burner is a $D-^3\text{He}$ tandem mirror (the Aegis/MetroVolt housings), a physically distinct machine with its own magnet: an end-plug field of **26.49 T** against the breeder's **16.84 T** centrepost. Its frozen design point is $Q_E = 1.318$, $P_{\text{fus}} = 4298.5 \text{ MW}$, $P_n = 233.95 \text{ MW}$, neutron fraction $f_n = 5.44\%$, ion-confining potential $\phi_i = 248.75 \text{ keV}$, central-cell $\beta_c = 0.55$ at $T_e = 89.72 \text{ keV}$, and a direct-energy-converter efficiency $\eta_{\text{DEC}} = 0.49$. The safety channel here is $s(x) = \beta_c$ with ceiling **0.60**; the toolkit β map gives $\beta_c(90 \text{ keV}) = 0.55$ (the design point), a barrier crossing at $T_i \approx 105.2 \text{ keV}$, and a local sensitivity $g(90) = 0.0046 \text{ keV}^{-1}$. Greedy heating drives β_c to **0.702** and violates the barrier on **95** of **100** steps; the certified clamp holds $\beta_c \leq 0.60$ with **0/100** violations, using command authority $u \in [82.2, 107.4] \text{ keV}$ without reaching the **70 keV** floor. Because Q_E falls with T_i above the design point ($1.318 \rightarrow 0.885$ at 120 keV), the β clamp also keeps the machine on the high side of the Q_E ridge (figure 3): the safety envelope and the performance envelope are aligned. The two clamps share the same projection 7 on two different plants and two different magnets — the architecture is machine-agnostic, the certificates are per-machine.



Certified backbone on the burner (KX-L1-A13-BU). The toolkit $\beta_c(T_i)$ map (blue) crosses the barrier $\beta_c = 0.60$ at $T_i \approx 105.2 \text{ keV}$; greedy heating reaches $\beta_c = 0.702$ at $\sim 128 \text{ keV}$ while the clamp holds $T_i \leq 105.2 \text{ keV}$. The engineering gain $Q_E(T_i)$ (rust, right axis) declines from **1.318** at the design point, so the β clamp keeps the machine on the good side of the Q_E ridge. The two curves are the toolkit β map and Q_E ridge sampled at the frozen anchor points.



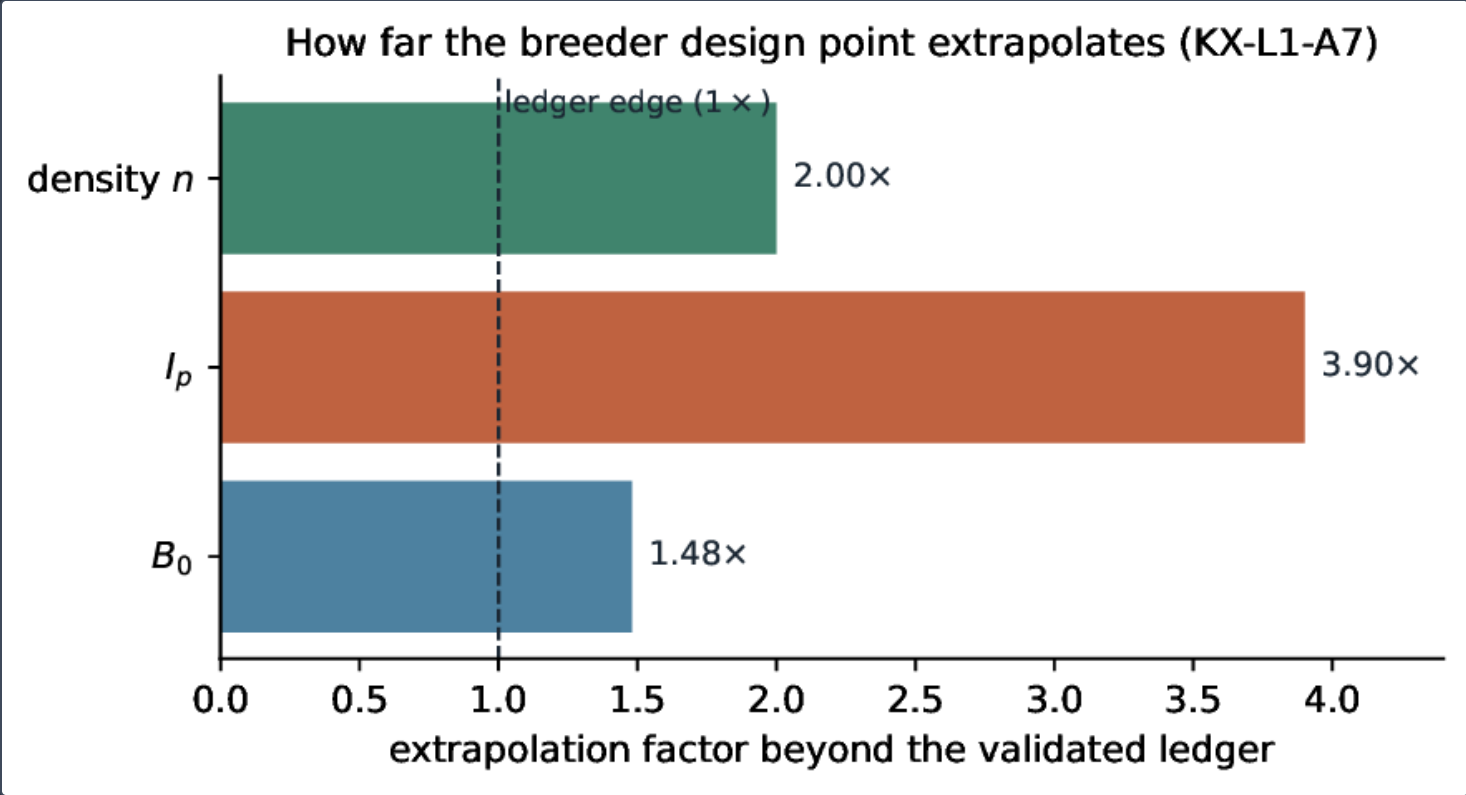
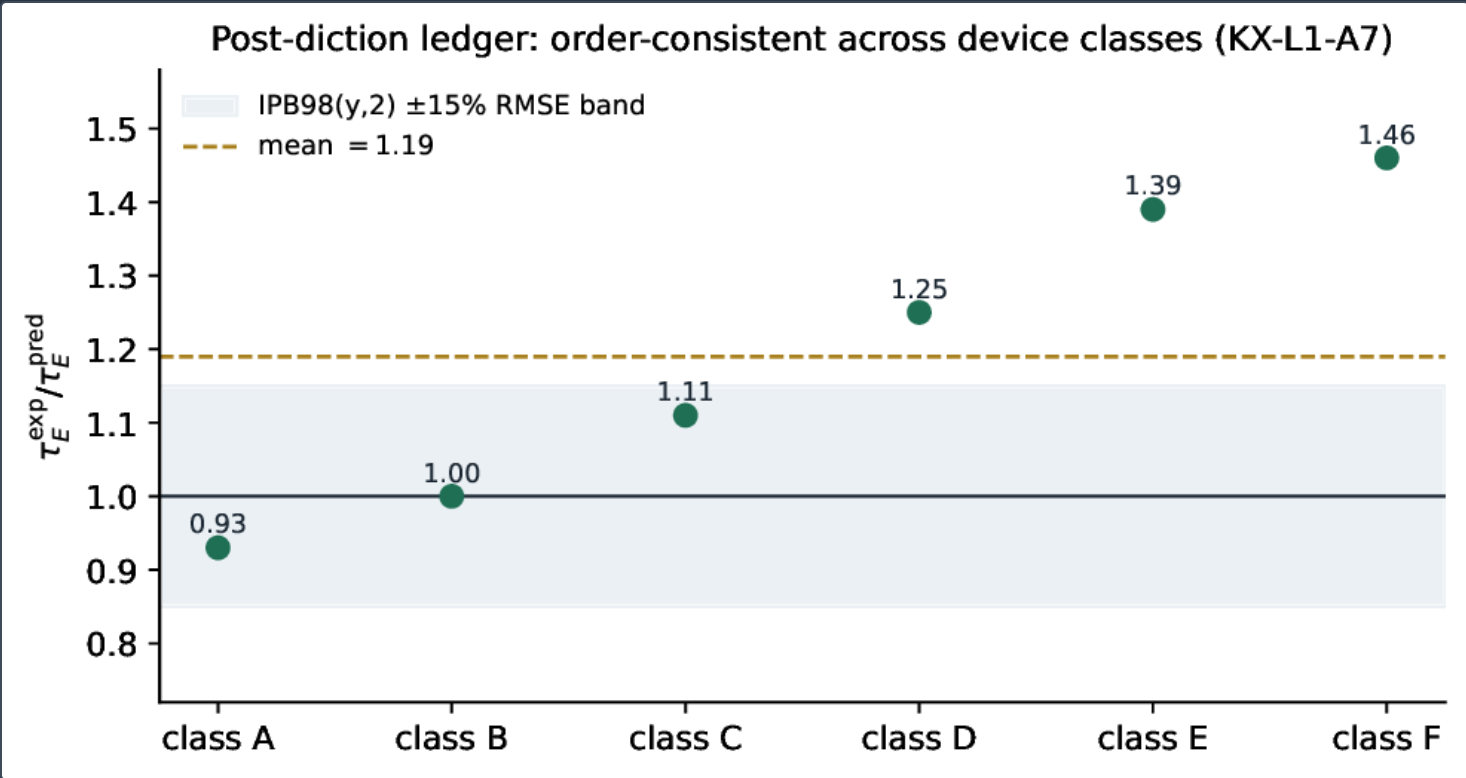
Demonstrated invariance on both magnets. Barrier violations per 100 steps fall from 50 to 0 on the breeder β_N limit (KX-L1-A13) and from 95 to 0 on the burner β_c limit (KX-L1-A13-BU). Two distinct machines, one certified projection.

QUANTITY	BREEDER (β_N , KX-L1-A13)	BURNER (β_c , KX-L1-A13-BU)
barrier ceiling $x_{\max}$	$\beta_N = 3.5$	$\beta_c = 0.60$
design-point anchor	$\beta_N = 1.532$ ($Q = 3.076$)	$\beta_c = 0.55$ ($Q_E = 1.318$)
uncertified peak	$\beta_N = 4.200$	$\beta_c = 0.702$
uncertified violations	50/100	95/100
clamped peak	$\beta_N = 3.500$	$\beta_c = 0.600$
clamped violations	0/100	0/100
barrier-margin minimum	$h = +0.000$	$h = +0.000$
authority used	$u \in [1.061, 1.250]$	$u \in [82.2, 107.4]$ keV
authority floor (unused)	$u_{\min} = 0.5$	$u_{\min} = 70$ keV

The certified backbone on two distinct magnets (frozen anchors KX-L1-A13, KX-L1-A13-BU). “Uncertified” is the greedy controller at maximum authority; “clamped” is the safety projection 7. The authority-used window and the barrier-margin minimum show the certificate is not authority-limited on either sequence.

THE EXTRAPOLATION LEDGER

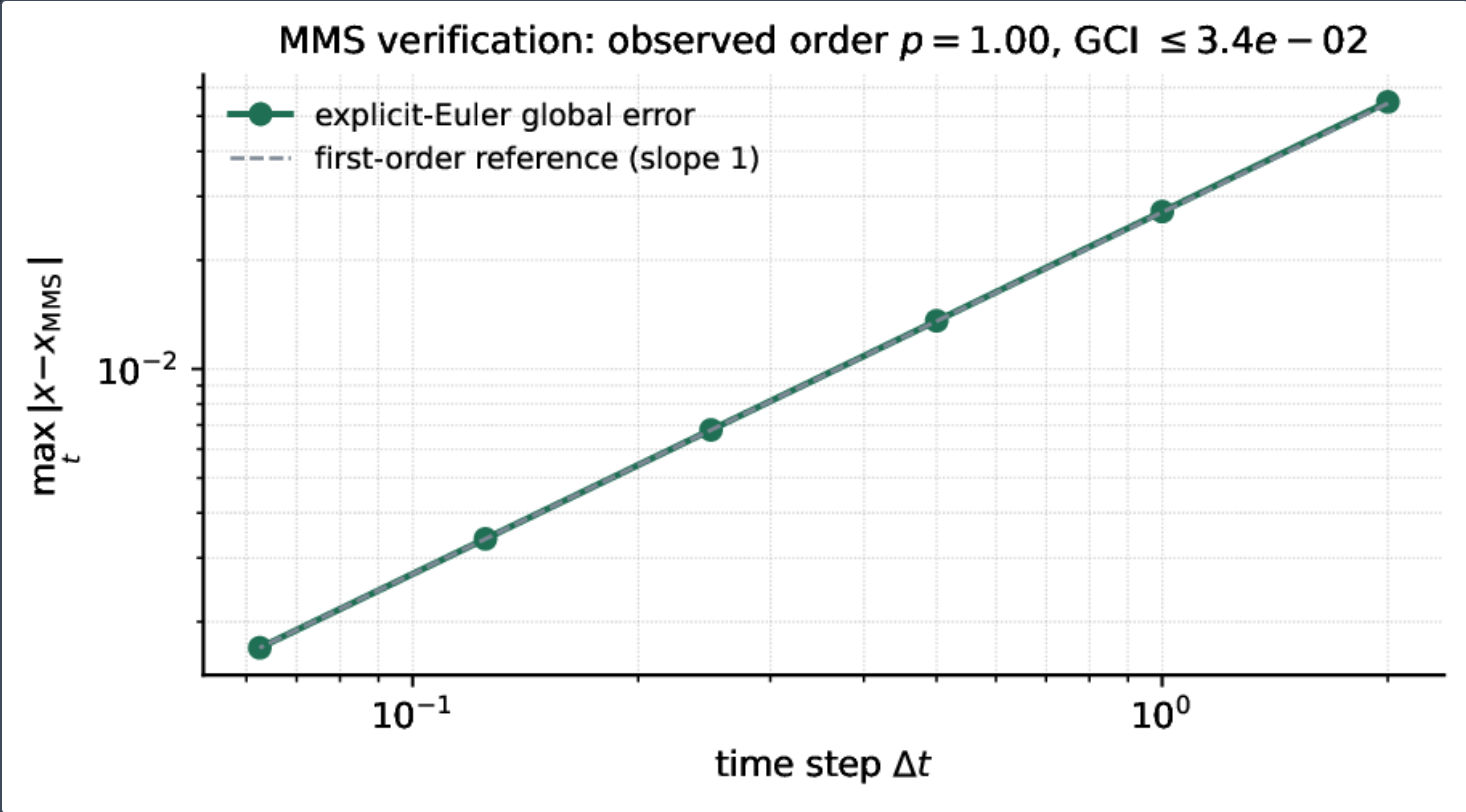
The extrapolation ledger (KX-L1-A7) bounds how far the underlying confinement model may be trusted. Across representative device classes the experiment-to-prediction confinement ratio 13 is $R_d = 1.19 \pm 0.21$ with range **0.93–1.46**, inside the IPB98(y,2) scaling’s own $\pm 15\%$ RMSE fit scatter (figure 5): the toolkit scaling is order-consistent with existing devices. The ledger then states the reach of the breeder design point honestly: at $\tau_E = 0.358 \text{ s}$ and $H_{98} = 1.0$ it extrapolates **1.48×** in B_0 , **3.9×** in I_p and **2.0×** in density beyond the ledger (figure 5). These reach factors ρ_l are exactly the coordinates that the epistemic gate 12 watches: where a query exceeds them, model authority is throttled to the certified floor. The ledger is the machine-readable statement of “here the models are on firm empirical ground, and here they are extrapolating.”



Post-diction ledger (KX-L1-A7). Ratio of experimental to predicted confinement time across representative device classes: 1.19 ± 0.21 , range 0.93–1.46, within the IPB98(y,2) $\pm$ 15% RMSE band (shaded). The scaling is order-consistent with existing devices.

NUMERICAL VERIFICATION

The clamp integrator's order is established by MMS (figure 6). Halving the step from $\Delta t = 2$ to 0.0625 drives the global error down at the first-order rate expected of explicit Euler, with observed order $p_{\text{obs}} = 1.00$ from 14 and a grid-convergence index bounded at the finest levels. The control results of table 1 are taken at a step well inside this asymptotic range, so the reported invariance is a property of the model and not of the discretisation.

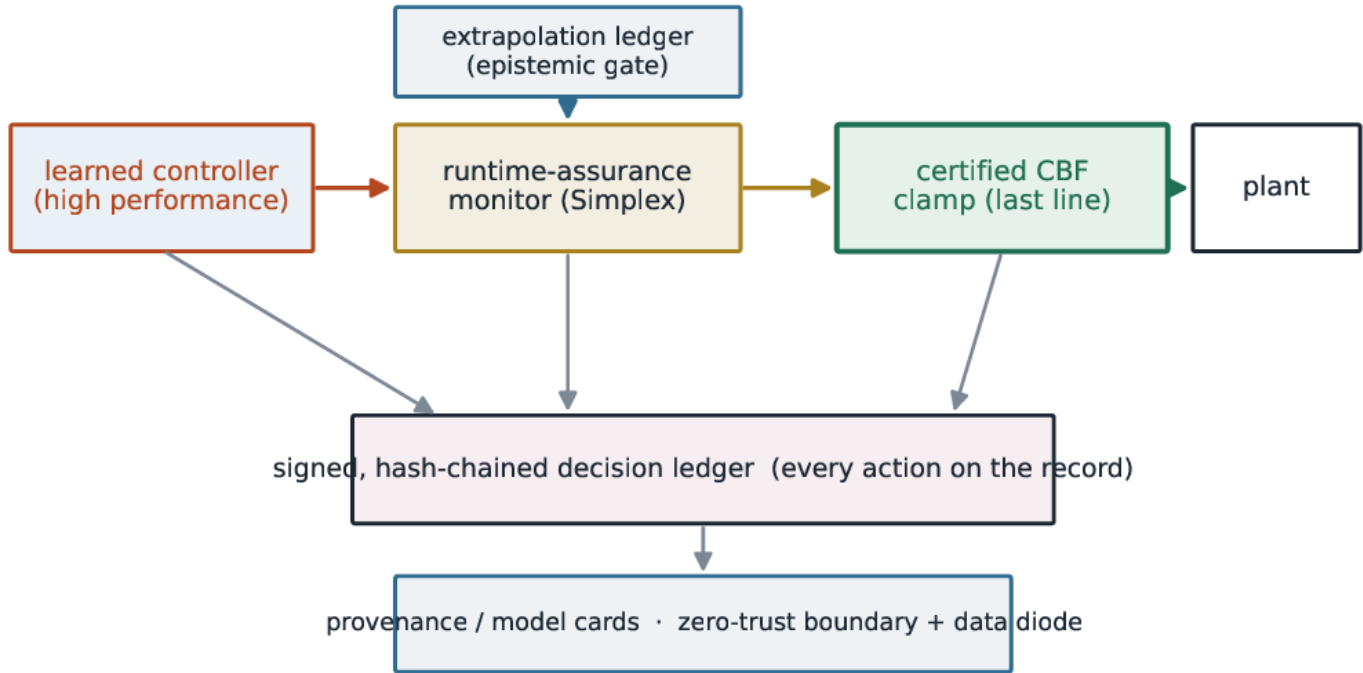


Method-of-manufactured-solutions verification of the clamp integrator. The explicit-Euler global error follows the first-order reference (slope 1); the observed order is $p_{\text{obs}} = 1.00$ and the GCI is bounded, confirming the control results are reported inside the asymptotic range.

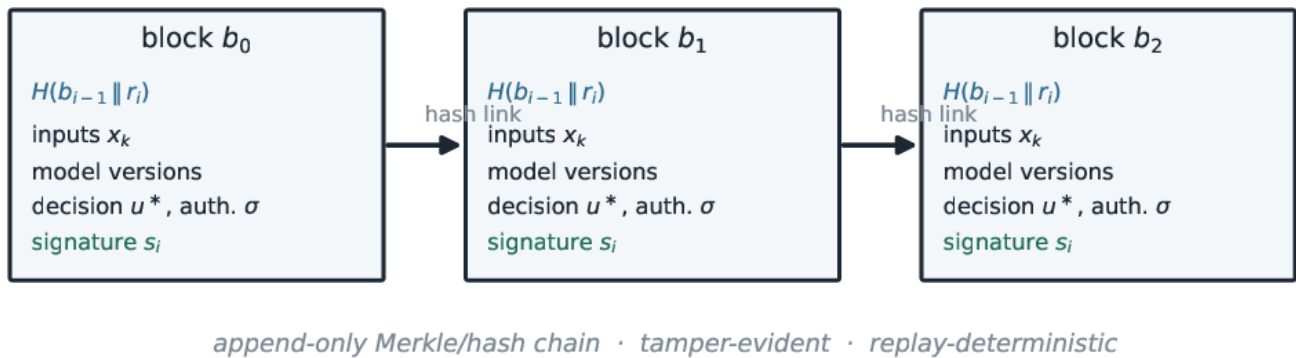
THE ASSURANCE LAYERS AND THEIR ACCEPTANCE CRITERIA

Around the demonstrated backbone the runtime-assurance monitor, signed ledger, zero-trust boundary and provenance records are specified against explicit, testable acceptance criteria (figure 7, table 2). The monitor of §2.5 arbitrates authority with the certified fallback; its criterion is that every switch and decision be justified and replay-auditable. The ledger of §2.6 chains every record (figure 8); its criterion is a clean end-to-end lineage audit. The gate of §2.7 blocks seeded boundary violations before they reach the plant. Figure 9 shows the authority arbitration as a state trajectory that hands control to the certified baseline the moment it nears the boundary. These are collected as inequalities and pass/fail criteria below.

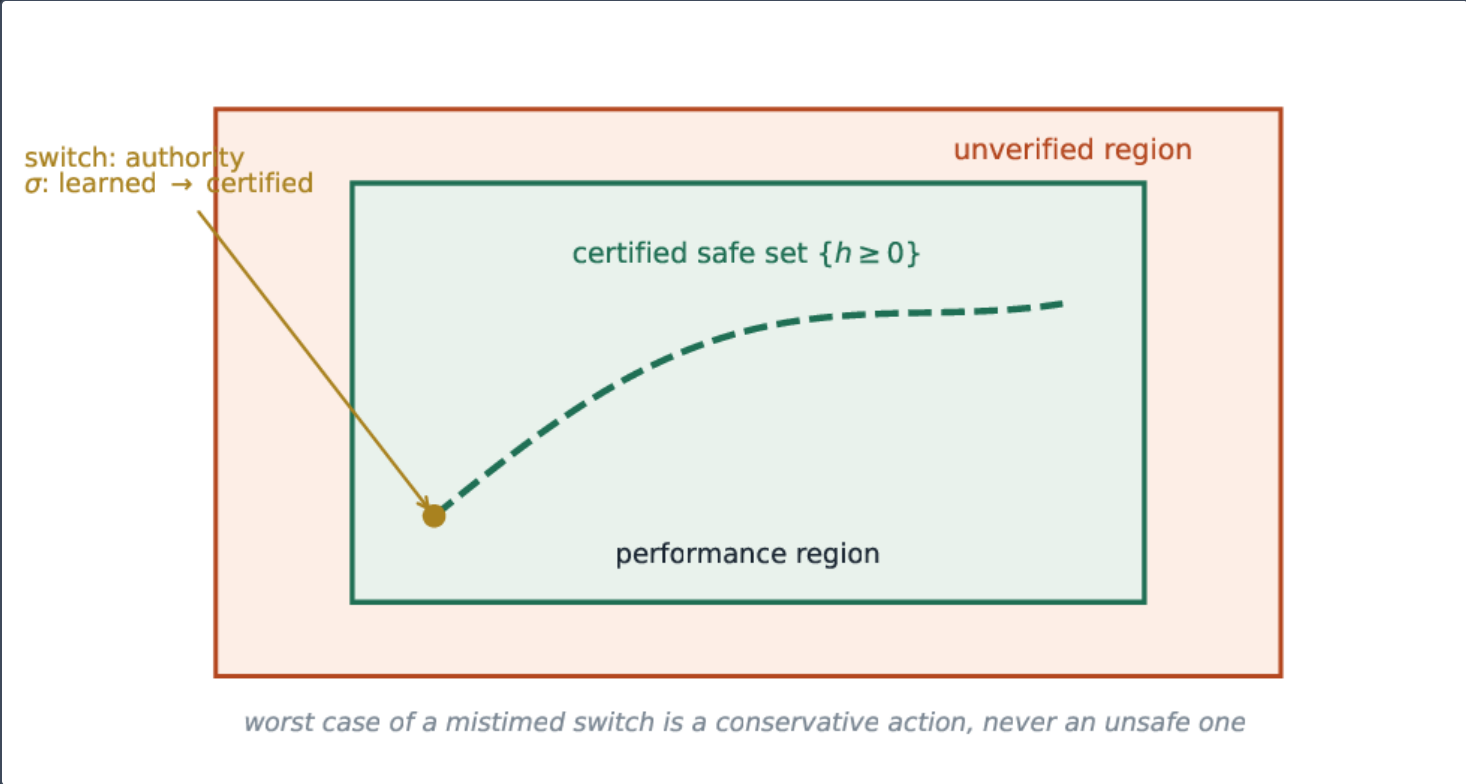
LAYER (SERIAL)	FUNCTION	RESULT / ACCEPTANCE CRITERION
certified CBF clamp (KX-L1-A13/-BU)	forward-invariant safe set	0/100 vs 50/100, 95/100 (demonstr.)
extrapolation ledger (KX-L1-A7)	bound out-of-range trust	$R_d = 1.19 \pm 0.21$ post-diction (demonstr.)
runtime-assurance monitor (SH-069)	Simplex authority switching	decision justified, replay-auditable
zero-trust monitor (SH-078)	block boundary violations	seeded violations blocked in sim
provenance / model cards (SH-077)	lineage on every model	lineage audit passes end to end
digital logbook (SH-063)	hash-chained action record	every action captured, hash-chained



(Schematic.) The last-line-of-defense architecture. The learned controller proposes; the Simplex-style runtime-assurance monitor arbitrates authority; the certified CBF clamp disposes before any command reaches the plant. An extrapolation ledger gates the learned controller epistemically; a signed, hash-chained decision ledger records every action; and provenance, model cards and a zero-trust boundary sit beneath the whole loop.



(Schematic.) The signed, hash-chained decision ledger of equations 10–11. Each block b_k carries the hash of its predecessor together with the inputs, the model versions, the applied decision and authority state, and a signature, forming an append-only Merkle/hash chain that is tamper-evident and replay-deterministic.



(Schematic.) Simplex authority arbitration of equations 8–9. The learned policy drives inside the performance region; as the state nears the certified safe-set boundary the decision module switches authority σ to the certified baseline (dashed). The worst case of a mistimed switch is a conservative action, never an unsafe one.

The acceptance criteria, collected as the contract that gates the architecture:

(A1) invariance	$h_k \geq 0 \ \forall k$ under (???)	(KX-L1-A13/-BU; demo.)
(A2) authority	$\bar{u}(x_k) \geq u_{\min} \ \forall k$	(reserve to spare; demo.)
(A3) arbitration	every decision justified, replay-auditable	(SH-069)
(A4) zero trust	seeded boundary violations blocked	(SH-078)
(A5) provenance	lineage audit clean end to end	(SH-077)
(A6) ledger	$b_k = H(b_{k-1} \ r_k)$ append-only, signed	(SH-063)
(A7) epistemic gate	$g_V(q) > 0 \Rightarrow \sigma \rightarrow C$	(KX-L1-A7)

Discussion

AGAINST THE SAFETY-FILTER LITERATURE

The certified clamp is a control-barrier safety filter, and its guarantee should be read against that literature ^[7,8,13]. Where CBF-QP filters are usually demonstrated on robotic or automotive plants, the contribution here is a filter bound to a specific frozen fusion design point and reported with measured invariance on the two machines to be built. The headline is not that a barrier can be enforced — that is theorem 4 — but that on the actual β_N and β_c limits, against disturbances that push an uncertified controller to $\beta_N = 4.200$ (50/100) and $\beta_c = 0.702$ (95/100), the projection holds 0/100 with authority to spare. The Simplex arbitration ^[12,4,14] and the reachability view of the switch ^[17] are the standard runtime-assurance machinery; our specific claim is that because the clamp filters in both authority modes, forward invariance is decoupled from the correctness of the switch — a mistimed switch costs performance, not safety.

AGAINST LEARNED FUSION CONTROL

Learned tokamak control has advanced quickly ^[1,2,3,5], and the value of the last-line-of-defense architecture is precisely that it lets those advances be deployed without betting safety on them. The learned controller of 9 may be any of these systems; the clamp only ever sees its proposed command u_{AI} and returns the nearest safe one. This is the same division of labour that classical plasma control ^[6] draws between performance and protection, made explicit and certifiable. The extrapolation ledger addresses the failure mode that learned systems add over classical ones — confident action out of distribution — by making the validated envelope a monitored object 12 rather than an implicit assumption.

AGAINST CONFINEMENT BENCHMARKS

The extrapolation ledger's post-diction places the breeder design point against existing devices honestly. A ratio of 1.19 ± 0.21 within the IPB98(y,2) $\pm 15\%$ fit scatter ^[23,24] means the toolkit scaling is neither optimistic nor pessimistic relative to the database; the reach factors ($1.48 \times B_0$, $3.9 \times I_p$, $2.0 \times$ density) then quantify where the design point genuinely extrapolates. Compact high-field concepts such as ARC and SPARC ^[34,35] and spherical-tokamak pilot plants ^[36] occupy the same extrapolation regime; stating it as a machine-readable ledger, rather than a narrative, is what lets the epistemic gate act on it in real time. The companion out-of-distribution paper (Kronos 2.79) develops the gate; the calibrated uncertainty that supplies the abstention threshold is developed in Kronos 2.77.

§ 6

Limitations

The certified backbone is demonstrated on a design-point-anchored one-dimensional reduced model of each safety channel; the barrier dynamics of β_N and β_c are surrogates that close the loop the static tools leave open, and the disturbance is a seeded perturbation rather than a validated plant excursion. The one open engineering item is to integrate the full runtime-assurance monitor and signed ledger on a hardware-in-the-loop bench and to re-establish the same acceptance criteria there. That integration composes the demonstrated clamp 7 and ledger 11 unchanged; it does not revise them.

Conclusion

We have formalised and demonstrated the Kronos last-line-of-defense architecture: a certified control-barrier safety projection at the base of an autonomous fusion controller, wrapped in Simplex-style runtime assurance, signed hash-chained provenance, a zero-trust boundary and a machine-readable extrapolation ledger. The certified core is a closed-form scalar clamp with a proven discrete-time invariance lemma, and it is demonstrated on two distinct magnets: $0/100$ barrier violations against a disturbance that drives an uncertified controller to $\beta_N = 4.200$ on $50/100$ steps on the breeder, and to $\beta_c = 0.702$ on $95/100$ steps on the burner, each with actuator authority to spare. The extrapolation ledger bounds trust ($R_d = 1.19 \pm 0.21$; reach $1.48 \times / 3.9 \times / 2.0 \times$), the runtime monitor arbitrates authority with a conservative fallback, and the signed ledger makes every decision auditable on replay. No learned component can cause an unsafe action, and every action the controller takes is on the record.

Acknowledgments Part of the Kronos Fusion Energy 2026 design series. The author thanks the Kronos control-and-assurance working group for the register anchors and the reproducibility discipline that underpins the demonstrated results.

Reproducibility The certified-clamp and extrapolation-ledger results are frozen anchors (KX-L1-A13, KX-L1-A13-BU, KX-L1-A7) in the Kronos Physics De-Risking Register [doi:10.5281/zenodo.22133057](https://doi.org/10.5281/zenodo.22133057) and are regenerable from the clamp trace and post-diction archives with the figure-generation script deposited alongside this paper (deterministic seed **20260726**). This paper is archived at [doi:10.5281/zenodo.22645817](https://doi.org/10.5281/zenodo.22645817) and its official version is maintained at <https://www.kronosfusionenergy.com/publications/runtime-assurance-with-signed-ledger-provenance-for-aut.html>. The formal-methods and runtime layers run on the in-house toolchain and FPGA/edge rig; the learned components the layers wrap are trained on the NVIDIA GPU HPC campaign.

Data availability This paper is archived at [doi:10.5281/zenodo.22645817](https://doi.org/10.5281/zenodo.22645817). The clamp traces, ledger records and post-diction tables are archived with the deposit and reference the Kronos Physics De-Risking Register [doi:10.5281/zenodo.22133057](https://doi.org/10.5281/zenodo.22133057) and the AI-and-quantum control paper (Kronos 1.5, [doi:10.5281/zenodo.22645702](https://doi.org/10.5281/zenodo.22645702), <https://www.kronosfusionenergy.com/publications/ai-quantum-control.html>). Any economic analysis is reported separately and is not included in the public deposit.

Competing interests Provisional patent applications relating to aspects of this work have been filed. The author declares no other competing interests.

A P P A R A T U S

Portfolio, People & Record

The intellectual-property estate, the people who built it, the limitations that gate it, and the sources it rests on.

1

GRANTED PATENT

4

2026 PROVISIONALS

40

TEAM MEMBERS

27

2022 PROVISIONALS

INTELLECTUAL PROPERTY

The patent portfolio

The estate spans a granted high-field magnet patent, pending utility and 2026 provisional filings that map to the two products, a digital-twin control provisional, a registered trademark application, and the original 2022 provisional family. Every patentable disclosure in the five-paper arXiv drop was protected **before** publication: the breeder and burner provisionals were filed 1–2 August, and a publication-gap omnibus filing the evening before the drop swept up the DEC, plasma-control, and REBCO-winding matter of the three otherwise-unprotected papers. Patent numbers and application serials are matters of public record; claim scope is summarised, not reproduced.

GRANTED & PENDING UTILITY

US 12,009,112	High-field tilted graded-REBCO magnet architecture (KRONO-002A) · app. 17/878,550	GRANTED
US 17/878,507	Core device / method (KRONO-001A) · utility	PENDING

2026 PROVISIONAL FILINGS — MAPPED TO THE PRODUCTS

64/124,209	Hyperion — spherical-tokamak tritium / helium-3 foundry · breeder · filed Aug 2026	FILED
64/124,220	Aegis / MetroVolt — synchrotron-cutoff, fuel-selection & plug-winding · generator · filed Aug 2026	FILED
64/115,167	KRONOS-CTRL digital-twin plant control · provisional · filed Jul 2026	FILED
64/005,440	Integrated spherical-tokamak fusion architecture · early integrated-architecture filing · filed Mar 2026	FILED
64/105,530	Negative-triangularity spherical-tokamak architecture · foundational ST filing · filed Jul 2026	FILED
64/128,097	Publication-gap omnibus — quasineutral two-species expander DEC · provenance-gated / latency-split plasma control · as-built high-field winding & conductor architecture · filed 7 Aug 2026, the evening before the arXiv drop	FILED

TRADEMARK & ORIGIN

FTK-98801894	Kronos Fusion Energy — trademark application	FILED
KR-2022-★	Original provisional family (KR-2022-00001 ... 00027) · 27 filings, 2022	PRIORITY

The narrow, defensible magnet novelty — the specific conductor and the digital-twin winding optimisation, not high-field REBCO as a category — is the commercial core that can earn ahead of any $Q > 1$ milestone, with markets in fusion magnet supply, MRI/NMR, accelerators, and proton therapy. The claim is scoped honestly: the high field is a *system field* result, not a stand-alone-coil record; the small-bore plug coil is structurally infeasible as a bare winding but resolved by a stress-managed structural shield (feasible-pending-FEA); and the winding-tape experiment returned a *null result*. The value is the method, not a field record.

THE TEAM

Founding partners, board, advisors & auditors

Kronos is built by a bench of advanced-fuel-fusion, high-field-magnet, direct-conversion, and materials specialists, with a board and operations team drawn from defense, national laboratories, and industry. Dates are shown as ranges; where a tenure has ended or is term-ending, the range says so.

FOUNDING PARTNERS — SCIENCE

Dr. Gerald Kulcinski

Co-Founder · Helium-3 & Advanced-Fuel Fusion
UW-Madison · 2023–Present

Dr. Carl Weggel

Chief Scientist / S.M.A.R.T. Design
MIT Alcator Program · 2022–Present

Dr. Robert J. Weggel

Magnetic Field Design
MIT Magnet Lab · 2022–Present

BOARD

Priyanca Ford

Founder & CEO · Executive Board
2022–Present

Bandel Carano

Finance / Strategy
Oak Investment Partners / Stanford · 2025–2026

Patrick Schweiger

Fusion Device Engineering
Oklo / CFS / TerraPower · 2025–Present

SCIENTIFIC ADVISORS

Dr. Steven O. Dean

Fusion Policy & History
DOE / Fusion Power Associates · 2025–Present

Dr. Patrick H. Diamond

Plasma Turbulence & Transport
UC San Diego · 2025–Present

Dr. Jack J. Dongarra

HPC & Numerical Algorithms
Turing Award · 2025–Present

Dr. Nasr M. Ghoniem

Chief Material Scientist
UCLA · 2025–Present

Dr. Siegfried Glenzer

Plasma Physics & Laser Fusion
Stanford / SLAC · 2022–Present

Dr. David A. Hammer

Pulsed-Power Fusion
Cornell · 2025–Present

Dr. Donald A. Spong

Plasma Theory & Confinement
ORNL · 2025–Present

Dr. Curtis Smith

Risk Assessment & Nuclear PRA
MIT / Idaho National Lab · 2025–Present

RADM (Ret.) David Goggins

Naval Engineering & Power-Plant Design
U.S. Navy · 2023–2026

Dr. Konstantin Batygin

Mathematician
Caltech · 2022–2025

Dr. Ruben Fair

Magnet Design / ITER Liaison
PPPL / Jefferson Lab · 2022–2025

Dr. Paul S. Weiss

Materials & Nanotechnology
UCLA · 2022–2025

SCIENTIFIC AUDITORS**Dr. Wilfred A. Cooper**

Plasma Equilibrium Theory
EPFL / Swiss Plasma Center · 2025–Present

Dr. Ahmed Hassanein

Plasma–Material Interactions
Purdue / Argonne · 2025–Present

Dr. Patrick E. Hopkins

Extreme Thermal Materials
U. of Virginia · 2025–Present

Dr. Peter Hosemann

Materials Joining & Structural Integrity
UC Berkeley · 2025–Present

Dr. Travis W. Knight

Advanced Nuclear Fuels & Systems
U. of South Carolina · 2025–Present

Dr. Philippe Lebrun

Cryogenics & Magnet Cooling
CERN · 2025–Present

Dr. Nitendra Singh

Nuclear Safety & Fuel Cycle
ITER · 2025–Present

Dr. Guido Van Oost

Magnetic Confinement & Fusion Systems
Ghent University · 2025–Present

Dr. Gary S. Was

Radiation Materials & Structural Integrity
U. of Michigan · 2025–Present

Dr. Ray Sedwick

Direct Energy Conversion
U. of Maryland · 2025

OPERATIONS**Michael Laughlin**

Chief Operating Officer
2022–Present

Martin Owens

Chief Strategy Officer
LANL / GE Hitachi · 2022–Present

MG (Ret.) Paul Pardew

Chief Contracting Officer
Army Contracting Command · 2022–Present

David Beck

Fusion Commercialization
U.S. Space Force · 2024–Present

Sushma Bhatia

Environmental & Legislative
Google · 2022–Present

Michael De Frenza

Technology Licensing
2023–Present

MG (Ret.) Robin L. Fontes

Cybersecurity & Defense
Army Cyber Command · 2022–Present

Vijay Gehani

Supply Chain & Components
INOX India · 2024–Present

Jon Michel Greenwood

IT & AI Infrastructure
Live Nation · 2023–Present

Brian C. O'Neill

National Security
CIA / ODNI · 2022–Present

Gen. (Ret.) Gustave F. Perna

DoD Liaison
Operation Warp Speed · 2022–2023

Andrea Romero

Marketing Lead
2022–Present

Legal counsel is retained; those roles are held on the internal roster and are not listed here.

SOURCES

References

- 1 J Degraeve, F Felici, J Buchli *et al.*, Magnetic control of tokamak plasmas through deep reinforcement learning, *Nature* **602**, 414–419 (2022). [doi:10.1038/s41586-021-04301-9](https://doi.org/10.1038/s41586-021-04301-9).
- 2 J Kates-Harbeck, A Svyatkovskiy and W Tang, Predicting disruptive instabilities in controlled fusion plasmas through deep learning, *Nature* **568**, 526–531 (2019). [doi:10.1038/s41586-019-1116-4](https://doi.org/10.1038/s41586-019-1116-4).
- 3 J Abbate, R Conlin and E Kolemen, Data-driven profile prediction for DIII-D, *Nucl. Fusion* **61**, 046027 (2021). [doi:10.1088/1741-4326/abe08d](https://doi.org/10.1088/1741-4326/abe08d).
- 4 J Seo, Y-S Na, B Kim *et al.*, Feedforward beta control in the KSTAR tokamak by deep reinforcement learning, *Nucl. Fusion* **61**, 106010 (2021). [doi:10.1088/1741-4326/ac121b](https://doi.org/10.1088/1741-4326/ac121b).
- 5 Y Fu, D Eldon, K Erickson *et al.*, Machine learning control for disruption and tearing mode avoidance, *Phys. Plasmas* **27**, 022501 (2020). [doi:10.1063/1.5125581](https://doi.org/10.1063/1.5125581).
- 6 M L Walker, D A Humphreys, D Mazon *et al.*, Fusion, tokamaks, and plasma control: an introduction and tutorial, *IEEE Control Syst. Mag.* **25**(5), 30–43 (2005). [doi:10.1109/MCS.2005.1512794](https://doi.org/10.1109/MCS.2005.1512794).
- 7 A D Ames, X Xu, J W Grizzle and P Tabuada, Control barrier function based quadratic programs for safety critical systems, *IEEE Trans. Autom. Control* **62**(8), 3861–3876 (2017). [doi:10.1109/TAC.2016.2638961](https://doi.org/10.1109/TAC.2016.2638961).
- 8 A D Ames, S Coogan, M Egerstedt, G Notomista, K Sreenath and P Tabuada, Control barrier functions: theory and applications, *2019 18th European Control Conference (ECC)*, 3420–3431 (2019). [doi:10.23919/ECC.2019.8796030](https://doi.org/10.23919/ECC.2019.8796030).
- 9 P Wieland and F Allgöwer, Constructive safety using control barrier functions, *IFAC Proc. Vol.* **40**(12), 462–467 (2007). [doi:10.3182/20070822-3-ZA-2920.00076](https://doi.org/10.3182/20070822-3-ZA-2920.00076).
- 10 F Blanchini, Set invariance in control, *Automatica* **35**(11), 1747–1767 (1999). [doi:10.1016/S0005-1098\(99\)00113-2](https://doi.org/10.1016/S0005-1098(99)00113-2).
- 11 S Prajna and A Jadbabaie, Safety verification of hybrid systems using barrier certificates, in *Hybrid Systems: Computation and Control*, LNCS **2993**, 477–492 (2004). [doi:10.1007/978-3-540-24743-2_32](https://doi.org/10.1007/978-3-540-24743-2_32).
- 12 L Sha, Using simplicity to control complexity, *IEEE Software* **18**(4), 20–28 (2001).
- 13 K L Hobbs, M L Mote, M C L Abate, S D Coogan and E M Feron, Runtime assurance for safety-critical systems: an introduction to safety filtering approaches for complex control systems, *IEEE Control Syst.* **43**(2), 28–65 (2023). [doi:10.1109/MCS.2023.3234380](https://doi.org/10.1109/MCS.2023.3234380).
- 14 J F Fisac, A K Akametalu, M N Zeilinger, S Kaynama, J Gillula and C J Tomlin, A general safety framework for learning-based control in uncertain robotic systems, *IEEE Trans. Autom. Control* **64**(7), 2737–2752 (2019). [doi:10.1109/TAC.2018.2876389](https://doi.org/10.1109/TAC.2018.2876389).
- 15 D Q Mayne, J B Rawlings, C V Rao and P O M Sokaert, Constrained model predictive control: stability and optimality, *Automatica* **36**(6), 789–814 (2000). [doi:10.1016/S0005-1098\(99\)00214-9](https://doi.org/10.1016/S0005-1098(99)00214-9).
- 16 E Garone, S Di Cairano and I Kolmanovsky, Reference and command governors for systems with constraints: a survey on theory and applications, *Automatica* **75**, 306–328 (2017). [doi:10.1016/j.automatica.2016.08.013](https://doi.org/10.1016/j.automatica.2016.08.013).
- 17 S Bansal, M Chen, S Herbert and C J Tomlin, Hamilton-Jacobi reachability: a brief overview and recent advances, *2017 IEEE 56th Conf. on Decision and Control (CDC)*, 2242–2253 (2017). [doi:10.1109/CDC.2017.8263977](https://doi.org/10.1109/CDC.2017.8263977).
- 18 R C Merkle, A digital signature based on a conventional encryption function, in *Advances in Cryptology — CRYPTO '87*, LNCS **293**, 369–378 (1988). [doi:10.1007/3-540-48184-2_32](https://doi.org/10.1007/3-540-48184-2_32).
- 19 S Haber and W S Stornetta, How to time-stamp a digital document, *J. Cryptology* **3**, 99–111 (1991). [doi:10.1007/BF00196791](https://doi.org/10.1007/BF00196791).
- 20 National Institute of Standards and Technology, *Secure Hash Standard (SHS)*, FIPS PUB 180-4 (2015). [doi:10.6028/NIST.FIPS.180-4](https://doi.org/10.6028/NIST.FIPS.180-4).
- 21 S Rose, O Borchert, S Mitchell and S Connelly, *Zero Trust Architecture*, NIST Special Publication 800-207 (2020). [doi:10.6028/NIST.SP.800-207](https://doi.org/10.6028/NIST.SP.800-207).
- 22 V Vovk, A Gammernan and G Shafer, *Algorithmic Learning in a Random World*, Springer (2005). [doi:10.1007/b106715](https://doi.org/10.1007/b106715).

- 23 ITER Physics Basis Editors *et al.*, Chapter 2: Plasma confinement and transport, *Nucl. Fusion* **39**(12), 2175–2249 (1999). [doi:10.1088/0029-5515/39/12/302](https://doi.org/10.1088/0029-5515/39/12/302).
- 24 E J Doyle *et al.* (ITPA), Chapter 2: Plasma confinement and transport, *Nucl. Fusion* **47**(6), S18–S127 (2007). [doi:10.1088/0029-5515/47/6/S02](https://doi.org/10.1088/0029-5515/47/6/S02).
- 25 F Troyon, R Gruber, H Saurenmann, S Semenzato and S Succi, MHD-limits to plasma confinement, *Plasma Phys. Control. Fusion* **26**(1A), 209–215 (1984). [doi:10.1088/0741-3335/26/1A/319](https://doi.org/10.1088/0741-3335/26/1A/319).
- 26 T C Hender *et al.*, Chapter 3: MHD stability, operational limits and disruptions, *Nucl. Fusion* **47**(6), S128–S202 (2007). [doi:10.1088/0029-5515/47/6/S03](https://doi.org/10.1088/0029-5515/47/6/S03).
- 27 P C de Vries *et al.*, Survey of disruption causes at JET, *Nucl. Fusion* **51**(5), 053018 (2011). [doi:10.1088/0029-5515/51/5/053018](https://doi.org/10.1088/0029-5515/51/5/053018).
- 28 L L Lao, H St John, R D Stambaugh, A G Kellman and W Pfeiffer, Reconstruction of current profile parameters and plasma shapes in tokamaks, *Nucl. Fusion* **25**(11), 1611–1622 (1985). [doi:10.1088/0029-5515/25/11/007](https://doi.org/10.1088/0029-5515/25/11/007).
- 29 J Candy, E A Belli and R V Bravenec, A high-accuracy Eulerian gyrokinetic solver for collisional plasmas, *J. Comput. Phys.* **324**, 73–93 (2016). [doi:10.1016/j.jcp.2016.07.039](https://doi.org/10.1016/j.jcp.2016.07.039).
- 30 P K Romano, N E Horelik, B R Herman, A G Nelson, B Forget and K Smith, OpenMC: a state-of-the-art Monte Carlo code for research and development, *Ann. Nucl. Energy* **82**, 90–97 (2015). [doi:10.1016/j.anucene.2014.07.048](https://doi.org/10.1016/j.anucene.2014.07.048).
- 31 P J Roache, Quantification of uncertainty in computational fluid dynamics, *Annu. Rev. Fluid Mech.* **29**, 123–160 (1997). [doi:10.1146/annurev.fluid.29.1.123](https://doi.org/10.1146/annurev.fluid.29.1.123).
- 32 W L Oberkampf and T G Trucano, Verification and validation in computational fluid dynamics, *Prog. Aerosp. Sci.* **38**(3), 209–272 (2002). [doi:10.1016/S0376-0421\(02\)00005-2](https://doi.org/10.1016/S0376-0421(02)00005-2).
- 33 C J Roy, Review of code and solution verification procedures for computational simulation, *J. Comput. Phys.* **205**(1), 131–156 (2005). [doi:10.1016/j.jcp.2004.10.036](https://doi.org/10.1016/j.jcp.2004.10.036).
- 34 B N Sorbom, J Ball, T R Palmer *et al.*, ARC: a compact, high-field, fusion nuclear science facility and demonstration power plant with demountable magnets, *Fusion Eng. Des.* **100**, 378–405 (2015). [doi:10.1016/j.fusengdes.2015.07.008](https://doi.org/10.1016/j.fusengdes.2015.07.008).
- 35 A J Creely *et al.*, Overview of the SPARC tokamak, *J. Plasma Phys.* **86**(5), 865860502 (2020). [doi:10.1017/S0022377820001257](https://doi.org/10.1017/S0022377820001257).
- 36 J E Menard *et al.*, Fusion nuclear science facilities and pilot plants based on the spherical tokamak, *Nucl. Fusion* **56**(10), 106023 (2016). [doi:10.1088/0029-5515/56/10/106023](https://doi.org/10.1088/0029-5515/56/10/106023).

COLOPHON

How this document was made

The Kronos Fleet — Combined Editorial, 2026 is set in **Fraunces** (display), **Gelasio** (text), and **IBM Plex Mono** (data & equations), carried forward from the Kronos editorial design system.

The figures are rendered at 300 dpi from the deposited generator scripts; the document is composed as a single self-contained file with embedded fonts and imagery, paginated to US Letter, and rendered to PDF through a headless Chromium engine in matched light and dark editions.

Every headline quantity carries an evidence class; withdrawn and restated values are marked as such. Nothing herein is frozen beyond the founder-approved Tier-A set.

Explore it live — turn the interactive [3-D model](#), run the deposited solvers in the [physics validator](#), and watch the [companion film series](#).



LEGAL NOTICE

Notice, status, and distribution

Conceptual design & simulation study. This document reports a conceptual design and simulation study. It is not a construction commitment, a safety-analysis report, a regulatory filing, or an offer of securities. Forward-looking statements — schedules, costs, market sizes, and performance — are estimates subject to the limitations set out in the Simulations part and may change.

Numbers and their classes. Quantities are reported with evidence classes and case labels; modelled, assumed, and requirement-class values are identified as such and are not to be quoted without their labels.

Public edition. This edition carries the physics and design only; all financial, funding, and defense-commercial content has been removed for public distribution. The scientific text corresponds to the arXiv and journal submissions.

© 2026 Kronos Fusion Energy. All rights reserved. Hyperion, Aegis, and MetroVolt are product designations of Kronos Fusion Energy.



THE 2026 KRONOS SERIES

One programme, many papers

This editorial is one paper in the Kronos 2026 series — a real fusion company's complete, reproducible physics record for a spherical-tokamak breeder and its low-neutron $D-^3\text{He}$ tandem-mirror burners. Every headline number is a frozen anchor in the Physics De-Risking Register.

Explore the whole programme: the [De-Risking Register](#), the interactive [3-D model](#), and the [physics validator](#).



KRONOS FUSION ENERGY

Runtime Assurance with Signed-Ledger Provenance for Autonomous Fusion Control: A Zero-Trust, Hardware- Failsafe Last Line of Defense Beneath the AI

A real fusion company building real machines. Explore the science, live.

[Zenodo · doi:10.5281/zenodo.22645817](https://doi.org/10.5281/zenodo.22645817)

[Read on kronosfusionenergy.com](https://kronosfusionenergy.com)

[The Physics De-Risking Register](#)