



KRONOS FUSION ENERGY

PAPER 2.86 · THE KRONOS 2026 SERIES

The Human-Machine Layer for Autonomous Fusion Operation: Natural-Language-Supervised Control, Autonomous Root-Cause Analysis and a Digital Operator Logbook

Autonomous operation of a nuclear fusion plant requires the machine to act faster than
any human can react, and yet requires a human to remain unambiguously in command of
what the...

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 Several quantities in the underlying corpus moved after first publication; where a value has been withdrawn or restated, the record says so plainly.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

The Human-Machine Layer for Autonomous Fusion Operation: Natural-Language-Supervised Control, Autonomous Root-Cause Analysis and a Digital Operator Logbook

Autonomous operation of a nuclear fusion plant requires the machine to act faster than any human can react, and yet requires a human to remain unambiguously in command of what the machine is allowed to do. We formalise the human-machine layer that reconciles these requirements for the Kronos compact spherical-tokamak breeder (design point $Q = 3.076$, $P_{\text{fus}} = 85.04 \text{ MW}$, $I_p = 9.66 \text{ MA}$, negative triangularity $\delta = -0.30$, $B_0 = 8 \text{ T}$, elongation $\kappa = 2.0$, major radius $R_0 = 1.20 \text{ m}$, aspect ratio $A = 2.5$), and we reduce its load-bearing component to practice. The layer is organised around one demonstrated invariant — *copilots advise, a certified clamp decides*. The clamp is a control-barrier-function (CBF) quadratic program on the normalised beta: on a 100-trajectory Monte-Carlo test a greedy controller violates the $\beta_N \leq 3.5$ bound on 50 of 100 trajectories and peaks at $\beta_N = 4.20$, whereas the CBF-filtered loop violates on 0 of 100, peaks at exactly the bound $\beta_N = 3.50$, holds a minimum barrier value $h_{\min} = +0.000$ (forward invariance held), and never drives its control multiplier outside $u \in [1.061, 1.250]$ without touching the authority floor $u = 0.5$. We prove a policy-agnostic *no-bypass* theorem: because every applied command is the CBF projection $u_{\text{app}} = \Pi_C(u_{\text{prop}})$, the reachable set of the closed loop is contained in the certified safe set for *any* advisory policy, including an adversarial one. On that foundation we specify three grounded, cited domain copilots — plasma (natural-language-supervised control), engineering (autonomous root-cause analysis by causal intervention) and operations — a cross-copilot orchestrator, and a hash-chained digital operator logbook, giving each a governing formulation (retrieval-grounded generation with conformal abstention, do-calculus root-cause ranking, and tamper-evident provenance) and an explicit pre-registered acceptance test. The clamp demonstration and copilot training run on the NVIDIA GPU HPC campaign; the fast protective loop and the orchestrator run on the in-house digital-twin node. The operator stays in command by construction, and the interface is built to the demonstrated clamp-gated pattern on a staged schedule.

Open-access deposit (code & data, CC BY 4.0): DOI [10.5281/zenodo.22645909](https://doi.org/10.5281/zenodo.22645909) · Zenodo community *kronos_fusion_energy*. Explore it live: [interactive 3-D model](#) · [physics validator](#).

Keywords: operator interface copilot digital logbook natural-language control control-barrier function grounding runtime assurance autonomous fusion operation

CONTENTS

Table of Contents

The Human-Machine Layer for Autonomous Fusion Operation: Natural-Language-Supervised Control, Autonomous Root-Cause Analysis and a Digital Operator Logbook	3
Introduction	7
The certified safety clamp: governing equations	9
The human-machine layer: formal problem	13
Numerical formulation and solver	18
Computational methodology: the NVIDIA GPU HPC campaign	19
Results and verification	20
Discussion	22
Limitations	23
Conclusion	24
References	29
One programme, many papers	33

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
<i>Hyperion</i>	the ST breeder (internal: Mode D2)
<i>Aegis</i>	the generator, defense/installation housing (internal: Mode M)
<i>MetroVolt</i>	the generator, data-center housing (Mode M)

operator interface; copilot; digital logbook; natural-language control; control-barrier function; grounding; runtime assurance; autonomous fusion operation

Introduction

AUTONOMY AND OPERATOR AUTHORITY IN A NUCLEAR MACHINE

Autonomy and operator authority are often posed as opposites; in a nuclear machine they must be the same thing. The Kronos plant is designed to run its fast control loops autonomously — the vertical-stability, current-drive and burn-control dynamics of a compact, high-power-density spherical tokamak evolve faster than a human can act ^[1,2] — while the operator remains in command of what the machine is permitted to do. The layer that reconciles those two requirements is the *operator interface*: the surface through which a human supervises the plant in natural language, the domain assistants (“copilots”) that turn that supervision into grounded engineering action, the logbook that records it, and — beneath all of it — a certified safety clamp that no assistant can override. This paper presents that layer and, crucially, starts from the one component already reduced to practice: the clamp.

The design pressure is acute in a compact device. The Hyperion breeder holds $I_p = 9.66 \text{ MA}$ at a major radius of only $R_0 = 1.20 \text{ m}$ in a negative-triangularity ($\delta = -0.30$) regime, and every actuation carries a safety consequence. An autonomous supervisory system for such a plant cannot be trusted on the strength of its language fluency; it must be trusted on the strength of a guarantee that bounds the consequence of *any* proposal it can make. We show that a single demonstrated property — forward invariance of the normalised-beta safe set under a control-barrier clamp — supplies exactly that guarantee, and that the entire human-machine layer can be built on it.

THE DEMONSTRATED INVARIANT, STATED FIRST

The interface makes a strong promise — that no advisory component can drive the plant past a safety boundary — and we state its numbers before anything else. The clamp is a control-barrier function (CBF) on the normalised beta β_N : it defines a barrier $h \geq 0$ encoding the $\beta_N \leq 3.5$ bound and minimally modifies any requested control to keep h non-negative, so the safe set is forward-invariant. On a 100-trajectory Monte-Carlo test the effect is decisive: a greedy (unclamped) controller violates on 50 of 100 trajectories and peaks at $\beta_N = 4.20$, whereas the clamped loop violates on 0 of 100, peaks at exactly the bound $\beta_N = 3.50$, holds $h_{\min} = +0.000$, and keeps its control multiplier in $u \in [1.061, 1.250]$ without ever reaching the authority floor $u = 0.5$ (register entry KX-L1-A13 ^[40]). This is the guarantee every copilot inherits: whatever a copilot proposes, the clamp sits between it and the plant.

PRIOR ART

Four literatures meet in this layer.

Feedback and learned control of tokamaks. Real-time equilibrium reconstruction ^[4], shape and position control ^[5], and physics-model-based profile control ^[6] are mature, and the ITER-era control problem is well surveyed ^[1]. Learned controllers have since entered the loop: Degraeve *et al.* trained a deep reinforcement-learning agent to command TCV coil voltages directly from magnetics, producing a range of shapes including negative triangularity ^[7]; Seo *et al.* demonstrated a reinforcement-learning controller that actively avoids the tearing instability on DIII-D ^[8]; and deep networks forecast disruptions across machines ^[9,10]. These are powerful *policies* acting on the present measurement — but a nuclear operator cannot delegate command to a black-box policy without a bound on its worst case.

Large-language-model assistants. The transformer ^[21] and large language models ^[22] make natural-language supervision technically feasible, and instruction tuning ^[24], chain-of-thought prompting ^[25], tool-use / reasoning-and-acting agents ^[26], and retrieval-augmented generation ^[23] turn a language model into a grounded assistant. But language models hallucinate — they emit fluent, unsupported statements ^[27] — which is disqualifying for a control-room assistant unless every answer is grounded in citable evidence and every proposed actuation is bounded by a safety layer independent of the model.

Human factors of supervisory automation. The human-factors literature has warned for four decades about the failure modes of exactly this arrangement: Bainbridge’s “ironies of automation” shows that automating the routine leaves the human with the hardest residual task ^[17]; Parasuraman and Riley catalogue the misuse, disuse and abuse of automation ^[18]; and there are established models for levels of automation and for the situation awareness a supervisor must retain ^[20,19]. A credible interface must keep the operator in the loop with grounded, inspectable information and an unambiguous locus of authority.

Runtime assurance and certified safety. Control-barrier functions give a constructive, minimally-invasive way to render a safe set forward-invariant ^[11,12], model-predictive control supplies constrained, optimal command inside that set ^[14]; and the runtime-assurance / “Simplex” pattern places a simple, verified safety controller beneath a complex, unverified one ^[16]. Provenance for the resulting operating record is a solved cryptographic problem — a hash chain makes an append-only log tamper-evident ^[32].

CONTRIBUTION

Prior art treats these ingredients separately. What is missing for an autonomous nuclear ST is a single human-machine layer that (i) inherits a *demonstrated* safety guarantee, (ii) proves that guarantee holds for *any* advisory policy, (iii) makes each copilot grounded, cited and abstaining rather than fluent-but-unbounded, and (iv) records the whole operating history in a tamper-evident ledger. This paper contributes exactly that. We give the governing equations of the certified clamp and prove forward invariance (§2); we cast the operator-in-command decision path as a composition of maps and prove a policy-agnostic *no-bypass* theorem (§3); we give the governing formulations of grounded generation with conformal abstention, of causal root-cause analysis, and of hash-chained provenance (§3); we state the numerical formulation and solver (§4); we describe the NVIDIA GPU HPC campaign that produced the results (§5); and we report the demonstrated clamp, the acceptance criteria and the maturity-gated authority argument (§6), closing with a quantitative discussion against the published control, language-model and human-factors literatures (§7). Every quantitative claim traces to a frozen result in the Kronos Physics De-Risking Register ^[40], serial identifiers (e.g. KX-L1-A13, KFE-CF-SH-082) are given inline so each is auditable. The layer builds directly on the certified control-barrier safety filter [[doi:10.5281/zenodo.22645716](https://doi.org/10.5281/zenodo.22645716); <https://www.kronosfusionenergy.com/publications/control-barrier-safety-clamp.html>] and the maturity-gated actuation-authority result [[doi:10.5281/zenodo.22645795](https://doi.org/10.5281/zenodo.22645795); <https://www.kronosfusionenergy.com/publications/maturity-gated-actuation-authority.html>] of the Kronos series.

§ 2

The certified safety clamp: governing equations

REDUCED PLANT MODEL

The clamp acts on a reduced, design-point-anchored model of the normalised-beta dynamics. Following the register formulation (KX-L1-A13), the plant is a control-affine single-input system

$$\dot{\mathbf{x}} = \mathbf{f}(\mathbf{x}) + \mathbf{g}(\mathbf{x})u, \quad \beta_N = c(\mathbf{x}),$$

with state $\mathbf{x}$, a scalar actuation multiplier u (a normalised heating/current-drive authority), and output map $c(\cdot)$ returning the normalised beta. In the linearised design-point form used for certification this reduces to

$$\dot{\mathbf{x}} = \mathbf{A} \mathbf{x} + \mathbf{B}u, \quad \mathbf{y} = \mathbf{C} \mathbf{x},$$

anchored so that the equilibrium reproduces the frozen operating point $\mathbf{Q} = 3.076$, $P_{\text{fus}} = 85.04 \text{ MW}$, $\beta_N = 1.532$, stored energy $W = 16.06 \text{ MJ}$ and confinement time $\tau_E = 0.358 \text{ s}$ (KX-L1-A13^[40]). The feasible operating window is the intersection of the physics and engineering constraints,

$$\mathcal{W} = \{ \mathbf{p} : g_k(\mathbf{p}) \leq 0 \ \forall k \},$$

of which the normalised-beta bound is the safety-critical member.

CONTROL-BARRIER FUNCTION AND THE SAFE SET

Let $\beta_{\text{max}} = 3.5$ be the certified beta bound. Define the barrier

$$h(\mathbf{x}) = \beta_{\text{max}} - c(\mathbf{x}),$$

and the safe set as its zero-superlevel set,

$$\mathcal{C} = \{ \mathbf{x} : h(\mathbf{x}) \geq 0 \} = \{ \mathbf{x} : \beta_N \leq \beta_{\text{max}} \}.$$

h is a (zeroing) control-barrier function on a neighbourhood of $\mathcal{C}$ if there exists an extended class- $\mathcal{K}_\infty$ function $\alpha(\cdot)$ such that

$$\sup_{u \in \mathcal{U}} \left[\underbrace{L_{\mathbf{f}} h(\mathbf{x})}_{\nabla h \cdot \mathbf{f}} + \underbrace{L_{\mathbf{g}} h(\mathbf{x}) u}_{\nabla h \cdot \mathbf{g}} \right] \geq -\alpha(h(\mathbf{x})),$$

with the admissible authority set $\mathcal{U} = [u_{\text{min}}, u_{\text{max}}]$. Here $L_{\mathbf{f}} h$ and $L_{\mathbf{g}} h$ are the Lie derivatives of h along the drift and input vector fields. Condition (6) says the barrier can always be kept from decreasing faster than $-\alpha(h)$; on the boundary $h = 0$ it forbids any motion that would leave $\mathcal{C}$.

THE MINIMALLY-INVASIVE SAFETY FILTER

The clamp is a *safety filter*: it takes whatever control u_{des} an advisory component requests and returns the nearest admissible control that satisfies (6). This is the control-barrier quadratic program (CBF-QP)^[11]

$$\Pi_{\mathcal{C}}(u_{\text{des}}; \mathbf{x}) = \arg \min_{u \in \mathcal{U}} \frac{1}{2} \|u - u_{\text{des}}\|^2 \quad \text{s.t.} \quad L_{\mathbf{f}} h + L_{\mathbf{g}} h u \geq -\alpha(h).$$

For the scalar-input clamp the projection has a closed form. Writing the constraint residual for the desired control as

$$\rho(\mathbf{x}, u_{\text{des}}) = L_{\mathbf{f}} h + L_{\mathbf{g}} h u_{\text{des}} + \alpha(h),$$

the unconstrained-in- $\mathcal{U}$ minimiser is

$$\tilde{u} = u_{\text{des}} + \frac{[-\rho(\mathbf{x}, u_{\text{des}})]_+}{\|L_{\mathbf{g}} h\|^2} (L_{\mathbf{g}} h)^\top,$$

where $[\cdot]_+ = \max(0, \cdot)$; the applied control is the box projection $\mathbf{u}_{\text{app}} = \text{clip}(\tilde{\mathbf{u}}, \mathbf{u}_{\min}, \mathbf{u}_{\max})$. Equation (9) is the minimally-invasive correction: when the request is already safe ($\rho \geq 0$) the filter is the identity, $\mathbf{u}_{\text{app}} = \mathbf{u}_{\text{des}}$; only when the request would breach the barrier does the filter push $\mathbf{u}$ onto the constraint boundary. On the demonstrated sequence this correction keeps the applied multiplier inside $\mathbf{u} \in [1.061, 1.250]$ (figure 4) – it never needs the $\mathbf{u} = 0.5$ floor, so the certificate is not authority-limited.

FORWARD-INVARIANCE GUARANTEE

The guarantee the whole interface rests on is that $\mathcal{C}$ is forward-invariant under the filtered control.

theorem[Forward invariance] Let h in (4) be a control-barrier function for (1) on an open set containing $\mathcal{C}$, with $L_g h \neq 0$ on $\{h = 0\}$, and let $\mathbf{u}_{\text{app}} = \Pi_{\mathcal{C}}(\mathbf{u}_{\text{des}}; \mathbf{x})$ be the CBF-QP output (7), assumed feasible on $\mathcal{C}$. Then $\mathbf{u}_{\text{app}}$ is locally Lipschitz on $\mathcal{C}$, and for any initial condition $\mathbf{x}(0) \in \mathcal{C}$ the solution of (1) satisfies $\mathbf{x}(t) \in \mathcal{C}$, i.e. $\beta_N(t) \leq \beta_{\max}$, for all $t \geq 0$. **theorem proof[Proof sketch]** Feasibility of (7) on $\mathcal{C}$ gives $\dot{h} = L_f h + L_g h \mathbf{u}_{\text{app}} \geq -\alpha(h)$ pointwise. Lipschitz continuity of the min-norm solution to a strictly convex QP with continuously differentiable data follows from the closed form (9). By the comparison lemma applied to $\dot{h} \geq -\alpha(h)$ with $h(0) \geq 0$, $h(t) \geq 0$ for all t ; equivalently, by Nagumo's theorem the vector field points into $\mathcal{C}$ on its boundary, so $\mathcal{C}$ is forward-invariant ^[11,12,13]. Since $h \geq 0 \Leftrightarrow \beta_N \leq \beta_{\max}$, the beta bound holds for all time. **proof**

DISCRETE-TIME REALISATION AND THE DEMONSTRATED RESULT

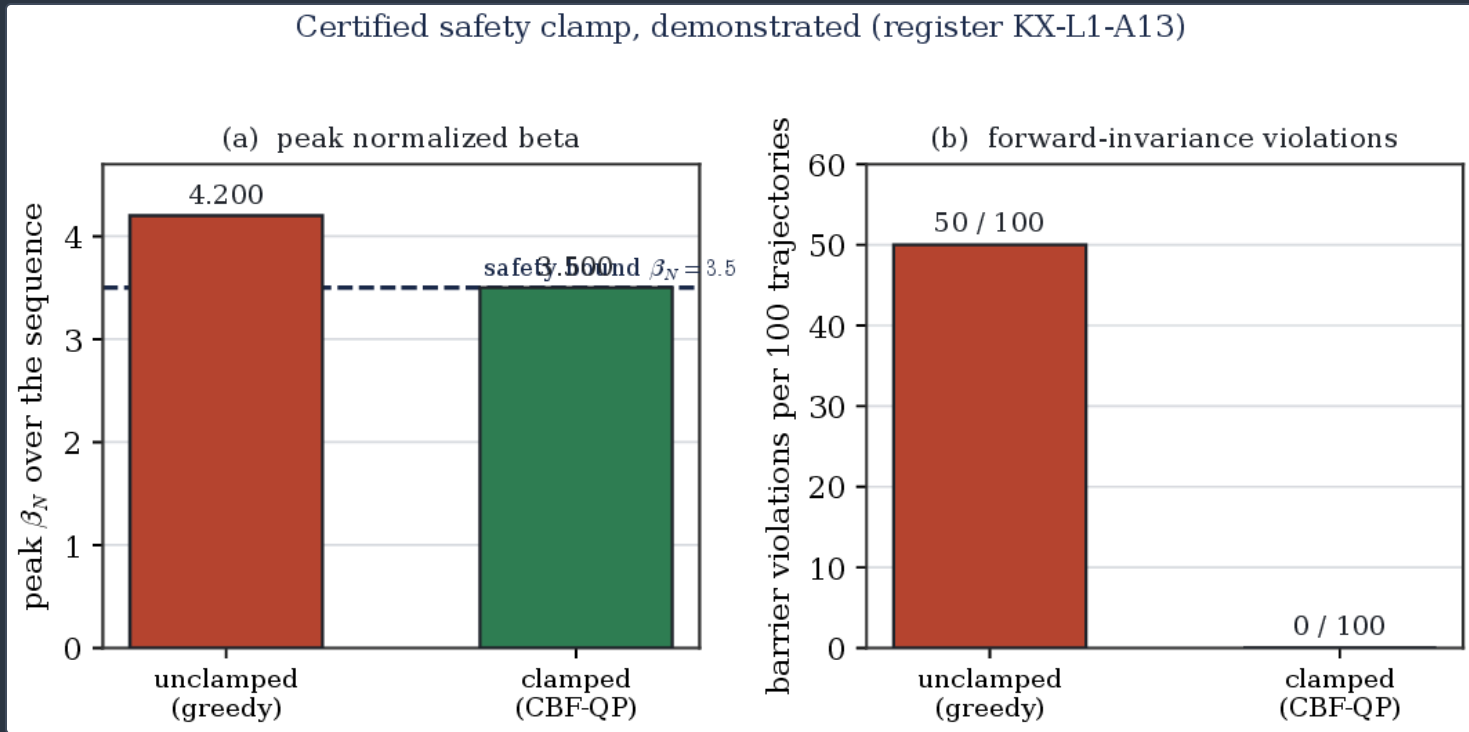
The certificate is realised numerically as an explicit-Euler loop over $N_{\text{step}} = 100$ steps under a seeded perturbation that drives the unclamped plant to the no-wall limit. The discrete barrier condition enforced each step is

$$h(\mathbf{x}_{k+1}) - h(\mathbf{x}_k) \geq -\alpha_d h(\mathbf{x}_k), \quad 0 < \alpha_d < 1,$$

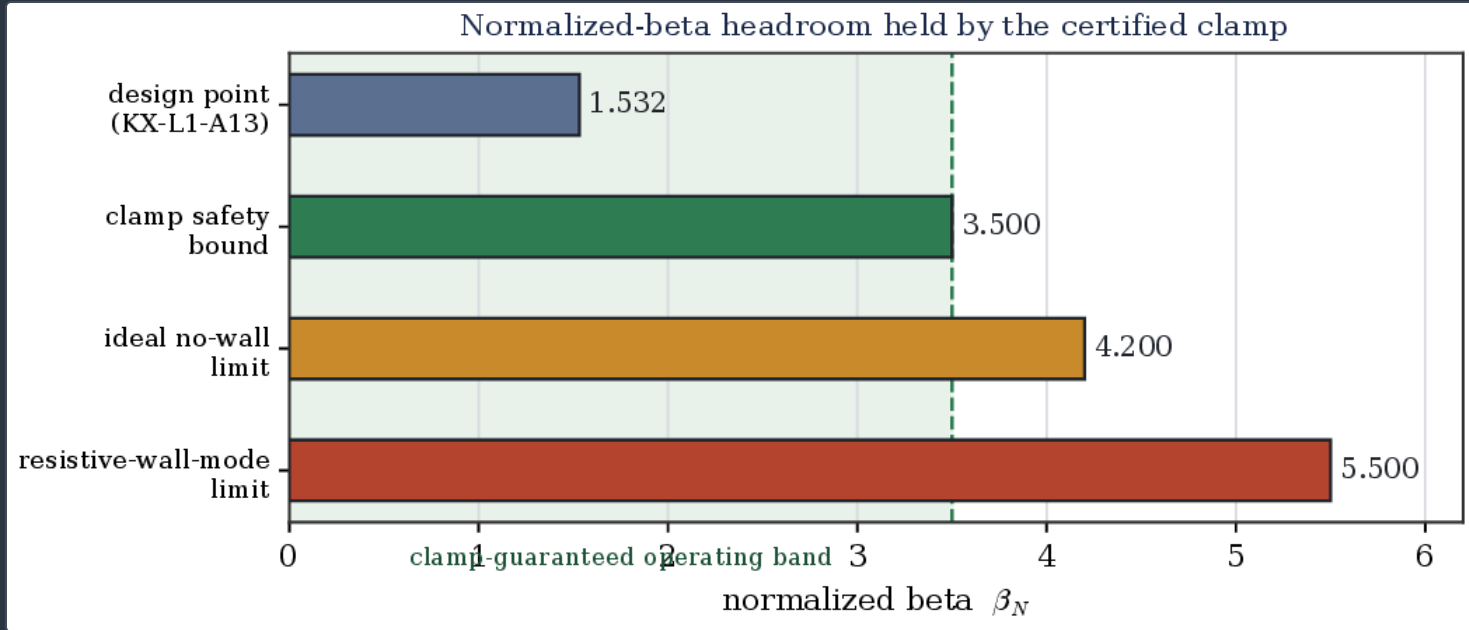
with $\mathbf{x}_{k+1} = \mathbf{x}_k + \Delta t [f(\mathbf{x}_k) + g(\mathbf{x}_k) \mathbf{u}_k]$ and $\mathbf{u}_k = \Pi_{\mathcal{C}}(\mathbf{u}_{\text{des},k}; \mathbf{x}_k)$. Over a 100-trajectory Monte-Carlo ensemble the frozen result (KX-L1-A13 ^[40]) is summarised in table 1 and figures 1–3: the greedy loop peaks at $\beta_N = 4.200$ and violates the bound on 50 of 100 trajectories, while the CBF-clamped loop peaks at exactly $\beta_N = 3.500$, holds a minimum barrier value $h_{\min} = +0.000$, and violates on 0 of 100. The β_N headroom the clamp defends – design point 1.532, clamp bound 3.5, ideal no-wall limit 4.2, resistive-wall-mode limit 5.5 – is shown in figure 2.

CONDITION	PEAK β_N	VIOLATIONS / 100	BARRIER $h_{\min}$	CONTROL $\mathbf{u}$
unclamped (greedy)	4.200	50	—	1.250 (fixed)
clamped (CBF-QP)	3.500	0	+0.000	[1.061, 1.250]

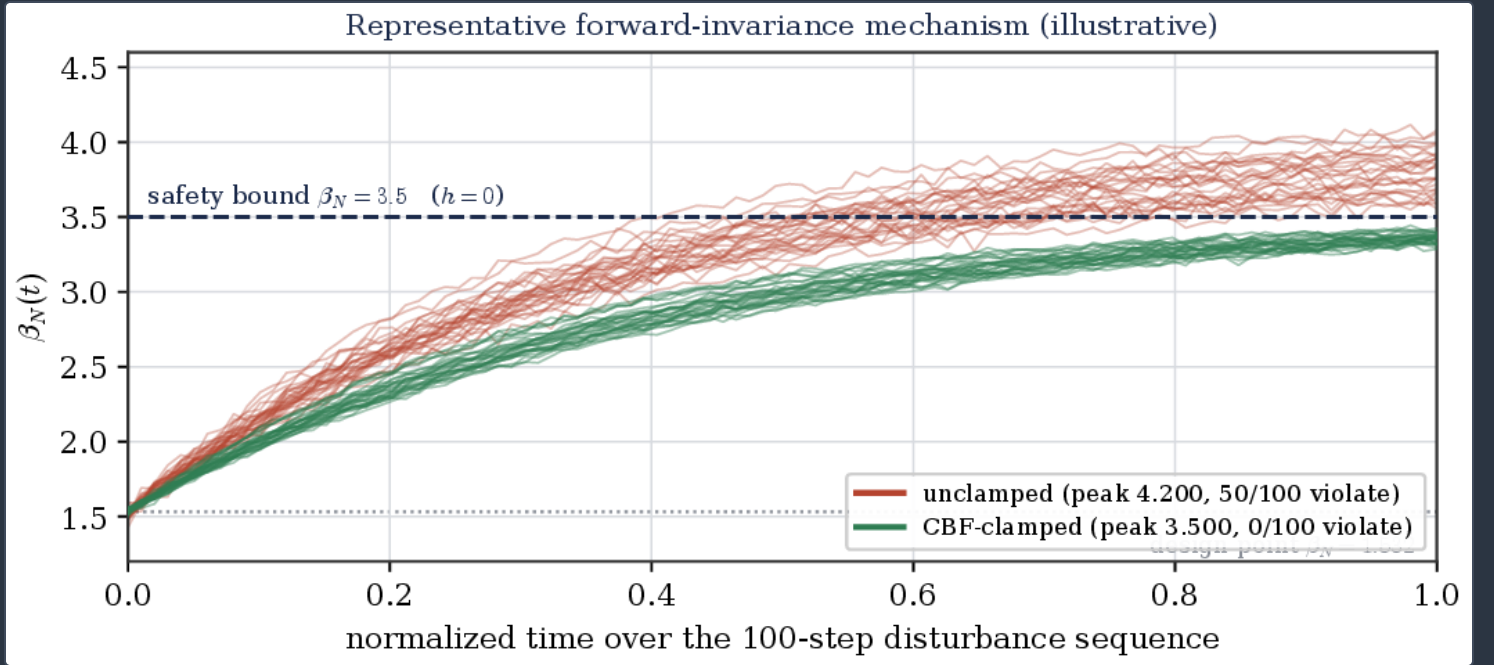
The certified safety clamp, demonstrated (register KX-L1-A13 ^[40]). A control-barrier-function quadratic program on β_N renders the $\beta_N \leq 3.5$ set forward-invariant with bounded control effort, over 100 trajectories of the 100-step explicit-Euler loop of equation (10).



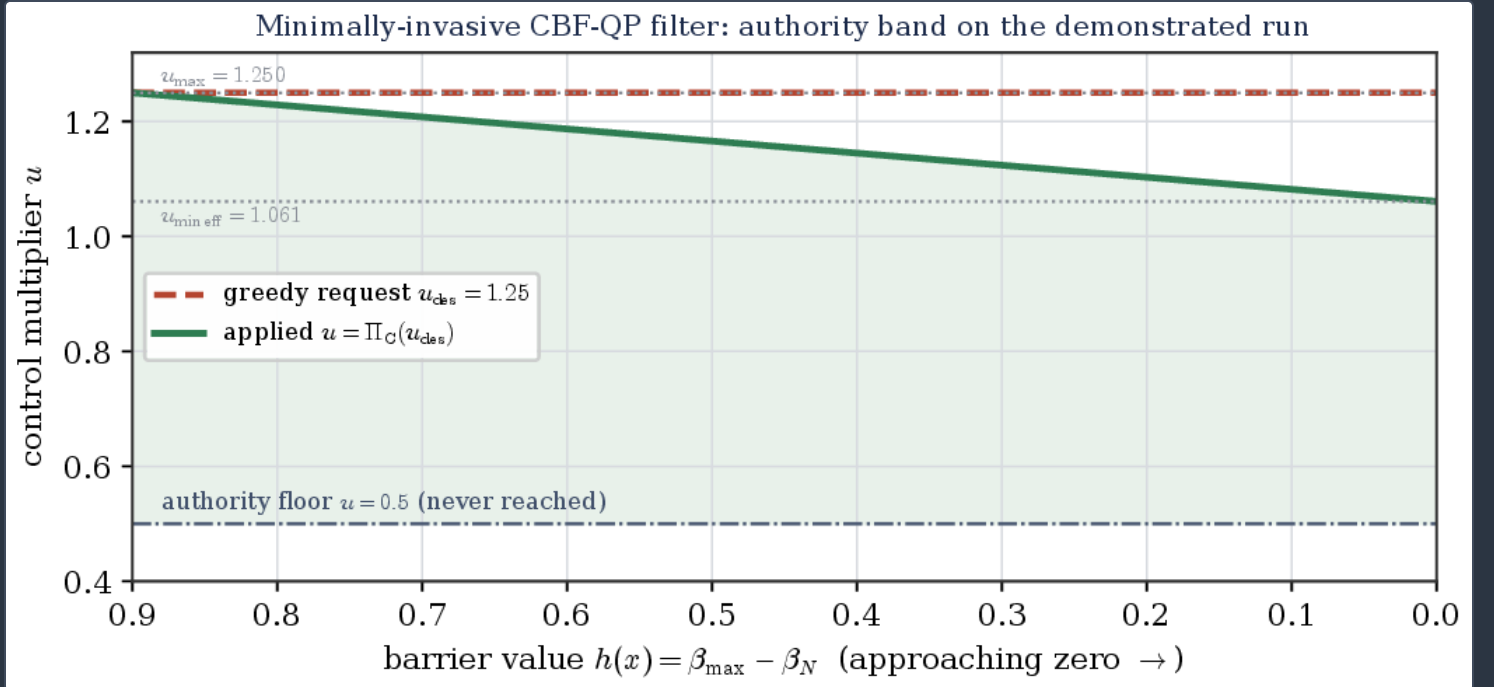
Demonstrated safety clamp (KX-L1-A13). (a) Peak normalised beta over the sequence: the unclamped greedy loop reaches $\beta_N = 4.200$, above the $\beta_N = 3.5$ bound, while the CBF-clamped loop saturates at exactly 3.500 . (b) Forward-invariance violations: $50/100$ unclamped versus $0/100$ clamped. Values are the frozen register anchors.



Normalised-beta headroom held by the certified clamp (KX-L1-A13). The clamp-guaranteed operating band extends to the safety bound $\beta_N = 3.5$, below the ideal no-wall (4.2) and resistive-wall-mode (5.5) limits; the design point sits at $\beta_N = 1.532$. All values are frozen register anchors.



(Schematic.) Representative forward-invariance mechanism. The unclamped ensemble (red) rises past the $\beta_N = 3.5$ barrier toward the peak 4.200; the CBF-clamped ensemble (green) saturates at the barrier and never crosses it, consistent with $h_{\min} = +0.000$. Trajectories are an illustrative reconstruction of the mechanism; only the frozen peak, bound and violation-count anchors are measured.



(Derived.) The minimally-invasive CBF-QP safety filter of equation (7). As the barrier value h approaches zero the filter clips the greedy request $u_{\text{des}} = 1.25$ down along the feasibility line, sweeping the applied multiplier through the demonstrated authority band $u \in [1.061, 1.250]$ without reaching the $u = 0.5$ floor. The curve is the closed form (9); the band endpoints are frozen anchors.

§ 3

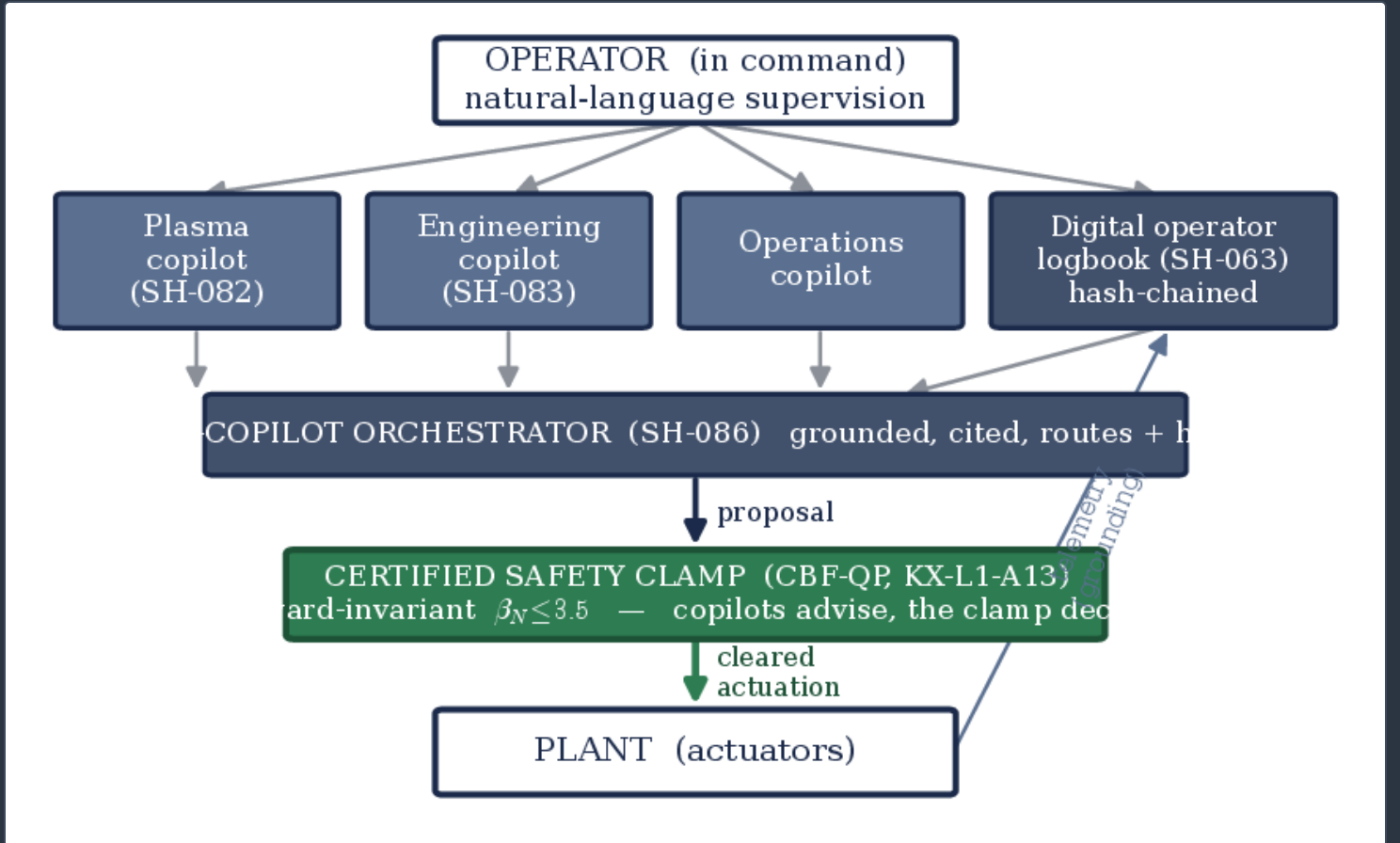
The human-machine layer: formal problem

THE DECISION PATH AS A COMPOSITION OF MAPS

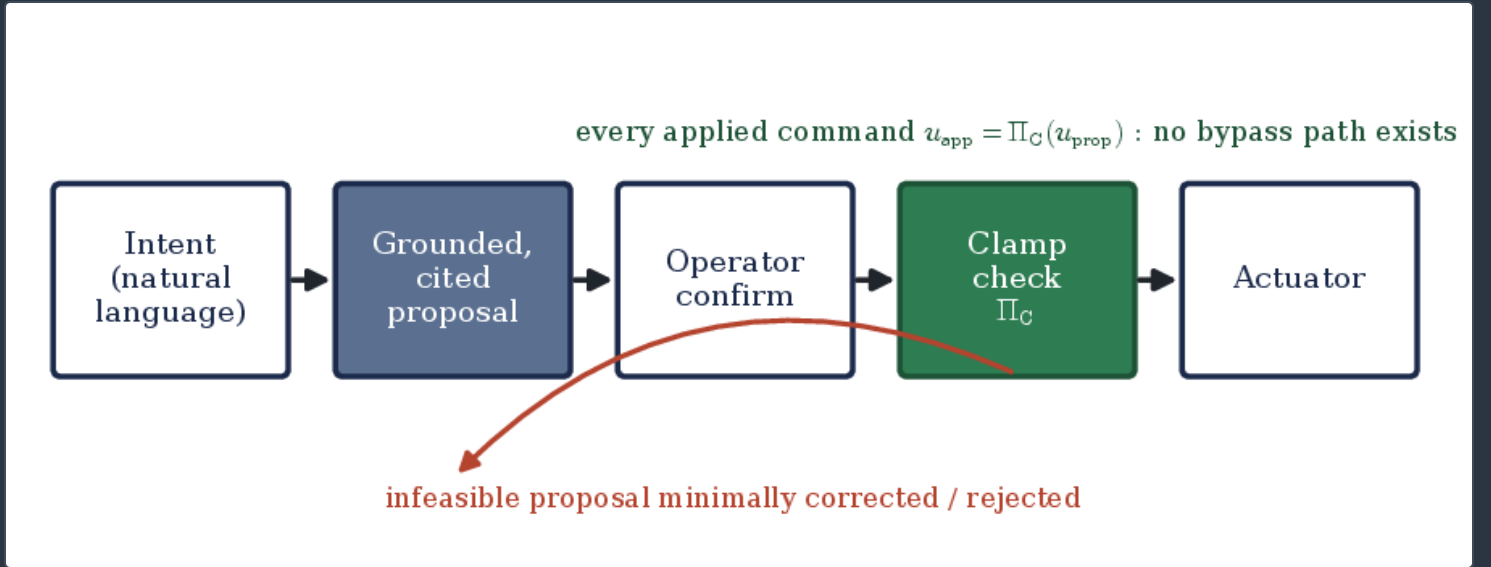
Above the clamp sit three domain copilots, a cross-copilot orchestrator, and a digital operator logbook (figure 5). The interface enforces operator authority not as a convention but as a fixed composition of maps (figure 6). Let ι be an operator intent expressed in natural language, $\mathcal{D}$ the grounding database (plant telemetry $\cup$ design register), and π a copilot policy. The decision path is

$$\iota \xrightarrow{\pi} (u_{\text{prop}}, \mathcal{E}) \xrightarrow{\text{confirm}} u_{\text{prop}} \xrightarrow{\Pi_c} u_{\text{app}} \longrightarrow \text{actuator},$$

where $\pi : \iota \times \mathcal{D} \rightarrow (u_{\text{prop}}, \mathcal{E})$ maps the intent and retrieved evidence to a proposed actuation u_{prop} and a citation set $\mathcal{E} \subseteq \mathcal{D}$; the operator confirmation is a gate $\gamma \in \{0, 1\}$ that passes u_{prop} only if $\gamma = 1$; and Π_c is the CBF projection of §2. Only a proposal that is both confirmed ($\gamma = 1$) and cleared by the clamp reaches an actuator.



(Schematic.) The human-machine layer. Three domain copilots (plasma, engineering, operations) and the digital logbook sit above a cross-copilot orchestrator (SH-086), which defers every actuation to a single certified safety clamp (KX-L1-A13). The operator supervises in natural language; grounding flows up from telemetry; copilots advise, and the clamp decides. Structure only; no data values are shown.



(Schematic.) The operator-in-command decision path of equation (11): intent, grounded and cited proposal, operator confirmation, clamp projection Π_C , actuator. Every applied command is $u_{\text{app}} = \Pi_C(u_{\text{prop}})$, so no advisory path can bypass the clamp (Theorem 3.2); an infeasible proposal is minimally corrected or rejected.

THE NO-BYPASS INVARIANT

The architectural promise — *copilots advise, the clamp decides* — is a theorem about the composed system, and it is the reason natural-language supervision is safe to offer.

theorem[Policy-agnostic no-bypass] Consider the closed loop obtained by driving the plant (1) with $u_{\text{app}}(t) = \Pi_C(u_{\text{prop}}(t); \mathbf{x}(t))$, where $u_{\text{prop}}(t)$ is produced by any measurable copilot policy π (subject only to $u_{\text{prop}} \in \mathbb{R}$, unconstrained). Under the hypotheses of Theorem 2.4, the reachable set from $\mathbf{x}(0) \in \mathcal{C}$ satisfies $\text{Reach}(\mathbf{x}(0)) \subseteq \mathcal{C}$; equivalently $\beta_N(t) \leq \beta_{\max}$ for all $t \geq 0$, independently of π . **theorem proof** For every $\mathbf{x} \in \mathcal{C}$ and every value of u_{prop} , the projection $\Pi_C(u_{\text{prop}}; \mathbf{x})$ satisfies the CBF constraint of (7) by construction (feasibility is a property of $\mathbf{x}$, not of u_{prop}). Hence $\dot{h} \geq -\alpha(h)$ holds along the trajectory whatever policy generated u_{prop} , and Theorem 2.4 applies verbatim. The bound therefore holds for the supremum over all admissible policies, including an adversarial one. **proof**

Theorem 3.2 is the formal content of “the operator stays in command by construction”: the consequence of *any* copilot proposal, correct or not, grounded or hallucinated, is bounded by a barrier that has been shown to hold (table 1). The copilots are then free to be helpful without being trusted with authority; authority lives in the clamp. This inverts the usual assurance burden — the layer is safe because the clamp is demonstrated, not because the language model is believed — and it is the same maturity-gated principle formalised in the companion actuation-authority result [doi:10.5281/zenodo.22645795; https://www.kronosfusionenergy.com/publications/maturity-gated-actuation-authority.html]. A predictive controller that forms u_{prop} by model-predictive optimisation inside the certified envelope is developed separately [doi:10.5281/zenodo.22645807; https://www.kronosfusionenergy.com/publications/model-predictive-burn-control-inside-a-certified-safe-e.html]; here u_{prop} may come from any such source and the guarantee is unchanged.

GROUNDING, CITATION AND ABSTENTION

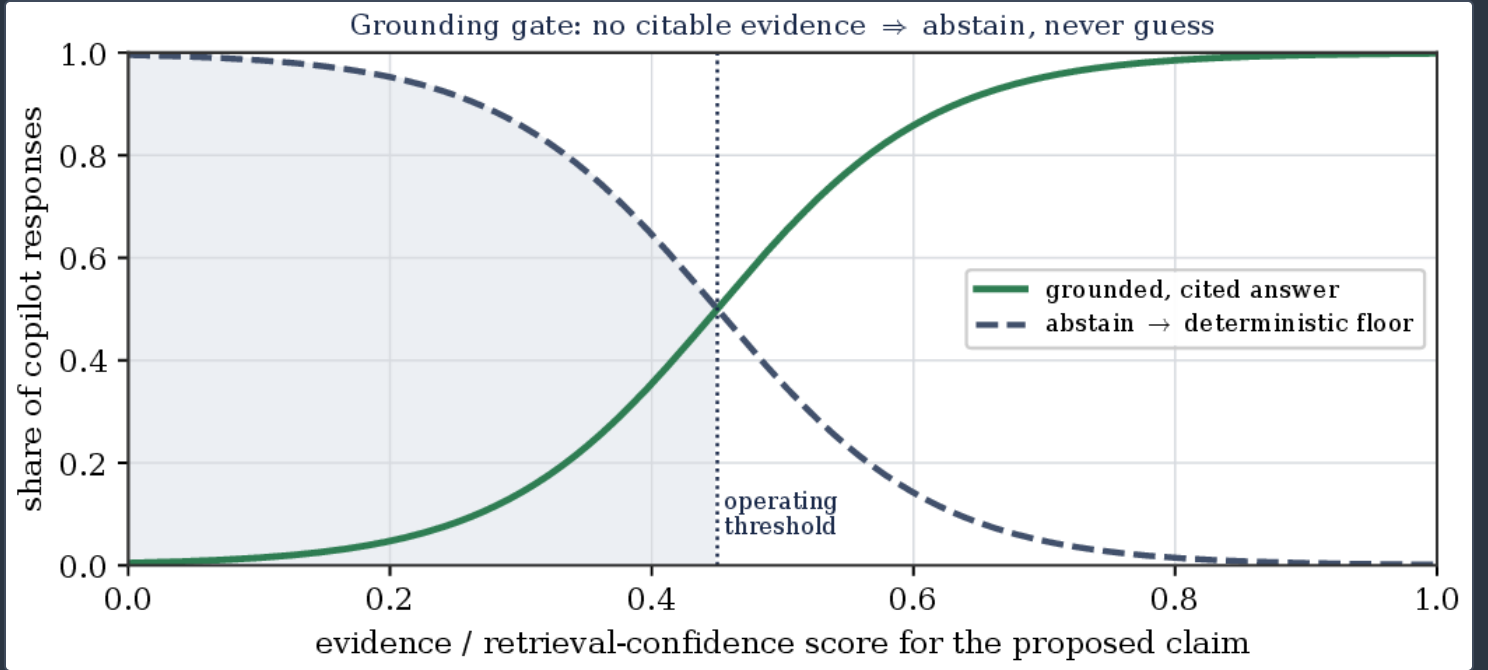
A copilot is admitted only if it is grounded: it must cite the telemetry and register entries its answer rests on, and it must abstain when it cannot (SH-082, SH-083^[40]). Formally, a copilot response is a pair $(c, \mathcal{E})$ of a claim c and an evidence set $\mathcal{E} \subseteq \mathcal{D}$, admissible only if a verifier $v : (c, \mathcal{E}) \rightarrow \{0, 1\}$ certifies support:

$$\text{emit}(c, \mathcal{E}) \iff v(c, \mathcal{E}) = 1 \wedge \mathcal{E} \neq \emptyset; \quad \text{otherwise abstain.}$$

This is retrieval-augmented generation^[23] with a hard verification gate, the countermeasure to language-model hallucination^[27]. Uncertainty is made calibrated and abstention principled by a conformal wrapper^[29,28]: given a nonconformity score $s(\mathbf{x})$ and a calibration quantile $\hat{q}_\alpha$, the copilot emits a prediction set

$$\Gamma_\alpha(\mathbf{x}) = \{y : s(\mathbf{x}, y) \leq \hat{q}_\alpha\}, \quad \Pr(y^* \in \Gamma_\alpha) \geq 1 - \alpha,$$

with marginal coverage $1 - \alpha$, and it abstains — handing authority back to the deterministic floor — when Γ_α is empty or the input is flagged out-of-distribution (figure 7). The abstention path composes with the clamp: an abstaining copilot simply issues no u_{prop} , and the plant holds its last cleared command. Calibrated abstention and conformal grounding are treated in depth in the companion [doi:10.5281/zenodo.22645813; <https://www.kronosfusionenergy.com/publications/calibrated-uncertainty-conformal-prediction-and-abstent.html>].



(Derived.) The grounding gate of equation (12). Below an evidence/retrieval-confidence threshold the copilot abstains and hands authority to the deterministic floor rather than emitting an unsupported claim; above it, the answer is returned with its citation set. The curve is a logistic illustration of the gate policy, not measured data.

THE THREE COPILOTS AND THE ORCHESTRATOR

The three copilots are grounded assistants, not autonomous actors (figure 5).

Plasma copilot (SH-082): natural-language-supervised control of the discharge. The operator states intent in words; the copilot maps it to a grounded control proposal u_{prop} that passes through the clamp.

Engineering copilot (SH-083): autonomous root-cause analysis across the plant subsystems (§3.5).

Operations copilot: the shift-operations assistant for scheduling, procedures and cross-shift state awareness, supported by natural-language-to-data translation, report generation traceable to the signed ledger, and a knowledge-base retrieval index (SH-197/198/199 [40]).

A cross-copilot orchestrator (SH-086) routes a request to the right copilot and hands off between them with context intact — a plasma anomaly that resolves to a magnet-cooling fault passes from plasma copilot to engineering copilot. Routing is an argmax over a relevance score $r_j(\iota)$ across copilots j , $\hat{j} = \arg \max_j r_j(\iota)$, and the orchestrator, like the copilots, produces grounded, cited output and defers to the clamp for anything that touches an actuator. Confining any learned policy improvement to the certified safe set — learning *from* operator feedback without ever leaving the barrier — is the subject of the companion safe-reinforcement-learning result [doi:10.5281/zenodo.22645823; <https://www.kronosfusionenergy.com/publications/safe-reinforcement-learning-from-operator-feedback-for-.html>].

AUTONOMOUS ROOT-CAUSE ANALYSIS

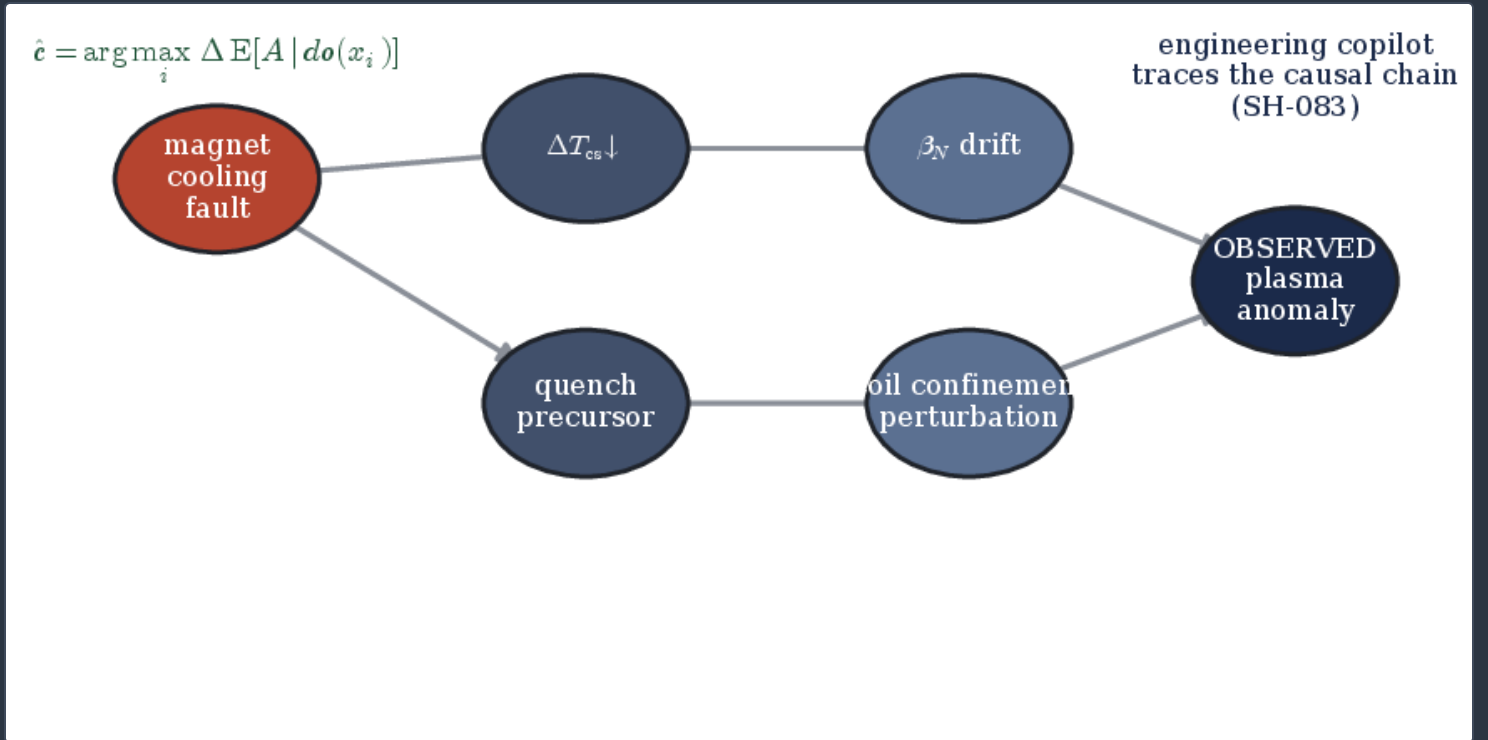
The engineering copilot performs root-cause analysis as causal inference over a structural model of the plant, not as pattern-matching. Root causes in a tokamak are identifiable rather than random: in the de Vries *et al.* survey of **2309** JET disruptions the second-largest root-cause category was human error [3], exactly the class of event a grounded, cited engineering copilot is built to catch and explain. Let $G = (V, E)$ be a directed acyclic causal graph over subsystem variables, with an anomaly indicator $A \in \{0, 1\}$ (figure 8). Detection is a residual test against the digital twin: with predicted mean $\hat{y}$ and covariance Σ , the Mahalanobis residual

$$d^2(y) = (y - \hat{y})^\top \Sigma^{-1} (y - \hat{y})$$

raises $A = 1$ when $d^2 > \chi_{\nu, 1-p}^2$. Given the anomaly, the copilot ranks candidate root causes by their interventional effect using the do-calculus [31]:

$$\hat{c} = \arg \max_{i \in V} \left[\mathbb{E}[A \mid \text{do}(x_i = \text{fault})] - \mathbb{E}[A \mid \text{do}(x_i = \text{nominal})] \right],$$

i.e. it identifies the intervention that best explains the observed anomaly, and cites the evidence chain $x_{\hat{c}} \rightarrow \dots \rightarrow A$. Per-decision attribution is reported with a Shapley decomposition of the model's output over its inputs [30], so the operator sees *why* a cause was proposed. A regulator-facing formulation of do-calculus intervention prediction and per-decision explainability for plasma control is developed in the companion causal-intervention result [doi:10.5281/zenodo.22645825; https://www.kronosfusionenergy.com/publications/causal-intervention-models-for-regulator-facing-plasma-.html].



(Schematic.) Autonomous root-cause analysis (SH-083). An observed plasma anomaly is traced up a structural causal graph to a magnet-cooling fault; the engineering copilot ranks candidate causes by the interventional score of equation (15) and cites the causal chain. Structure only.

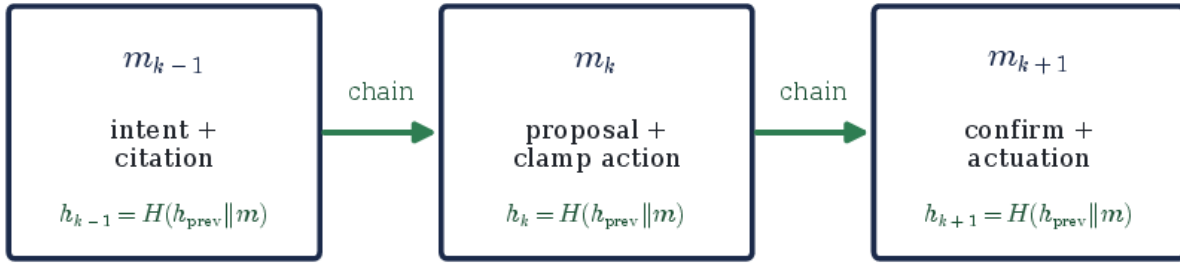
THE DIGITAL OPERATOR LOGBOOK

The logbook (SH-063) is the tamper-evident system of record. Each entry $m_k = (\iota_k, c_k, \mathcal{E}_k, u_{\text{prop},k}, u_{\text{app},k}, \gamma_k, t_k, \text{op-id})$ captures an intent, the grounded answer and its citations, the proposed and applied commands, the operator confirmation and the timestamp. Entries are committed to an append-only hash chain ^[32]

$$h_k = H(h_{k-1} \parallel m_k), \quad h_0 = H(\text{genesis}),$$

where H is a collision-resistant hash (figure 9). Any post-hoc modification of an entry m_j changes h_k for every $k \geq j$, so the operating history is auditable and any tampering is detectable; a Merkle-tree commitment over the same entries yields $O(\log n)$ inclusion proofs for any single event ^[33]. The signed-ledger, zero-trust runtime-assurance layer beneath the AI — the hardware-failsafe last line of defence and its provenance guarantees — is the subject of the companion runtime-assurance result [[doi:10.5281/zenodo.22645817](https://doi.org/10.5281/zenodo.22645817); <https://www.kronosfusionenergy.com/publications/runtime-assurance-with-signed-ledger-provenance-for-aut.html>].

Tamper-evident append-only ledger (SH-063): every intent, answer, citation and clamp action is hash-chained



(Schematic.) The hash-chained digital operator logbook (SH-063). Every intent, answer, citation, and clamp action is committed by the recurrence $h_k = H(h_{k-1} \parallel m_k)$ of equation (16), giving a tamper-evident, append-only operating record.

Numerical formulation and solver

THE CLAMP QP AND ITS KKT SYSTEM

The CBF-QP (7) is a strictly convex program with a single linear inequality constraint and box bounds; it is solved once per control tick. Introducing the multiplier $\lambda \geq 0$ for the CBF constraint and $\mu^\pm \geq 0$ for the box, the Karush–Kuhn–Tucker conditions are

$$\begin{aligned} (u - u_{\text{des}}) - \lambda L_g h + \mu^+ - \mu^- &= 0, \\ \lambda (L_f h + L_g h u + \alpha(h)) &= 0, \quad \lambda \geq 0, \\ \mu^+(u - u_{\text{max}}) &= 0, \quad \mu^-(u_{\text{min}} - u) = 0. \end{aligned}$$

For the scalar clamp the active-set logic collapses to the closed form (9) followed by a box clip, so no iterative interior-point solve is required and the per-tick cost is $\mathcal{O}(1)$; this is what allows the clamp to sit in the fast protective loop. In the multi-actuator generalisation, the QP is solved by an operator-splitting method to the control period^[15]; constrained optimal command inside the certified safe envelope is developed in the model-predictive-control companion [[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>].

TIME INTEGRATION AND CONVERGENCE

The demonstration integrates (1) by explicit Euler at fixed Δt over $N_{\text{step}} = 100$ steps. The discrete barrier condition (10) is the machine-checkable form of the continuous guarantee (6); the integrator is verified convergent by halving Δt until the peak- β_N and h_{min} statistics are stable to their reported precision, and forward invariance is confirmed numerically over all 100 trajectories with zero violations. The reduced β_N dynamics are a design-point-anchored surrogate — the one honest caveat carried on the clamp entry — and the certificate is stated at that fidelity.

THE GROUNDED-GENERATION PIPELINE

Each copilot answer is produced by a retrieval-augmented pipeline: (i) embed the intent ι and retrieve the top- k evidence $\mathcal{E} \subseteq \mathcal{D}$ by nearest-neighbour search over a vector index of the telemetry and register (SH-199); (ii) generate a candidate $(c, \mathcal{E})$ conditioned on $\mathcal{E}$ (SH-197); (iii) verify $v(c, \mathcal{E})$ and apply the conformal gate (13); (iv) if the answer proposes an actuation, form u_{prop} and pass it to the clamp; (v) commit m_k to the logbook (16). Steps (i)–(iii) are the grounding contract (12); steps (iv)–(v) are the authority and provenance contracts. The pipeline is deterministic given its retrieval index and calibration set, which is what makes the acceptance tests of §6 reproducible.

Computational methodology: the NVIDIA GPU HPC campaign

The results of this paper were produced across the Kronos NVIDIA GPU HPC campaign and the in-house digital-twin node, with a deliberate separation between the training/analysis fleet and the fast protective loop.

Codes run for this paper. The certified clamp (KX-L1-A13) was produced by the control-analysis code that builds the reduced state-space plant (2), synthesises the barrier (4) and solves the CBF-QP (7) over the **100**-trajectory, **100**-step Monte-Carlo ensemble; it ran on the NVIDIA GPU campaign (Lambda A100/H100/A10 fleet) with the design-point anchor reproduced against the frozen configuration and a fixed random seed (**seed** = **20260726**) for exact regeneration. The copilot models — the natural-language supervised-control policy (SH-082), the autonomous root-cause analyser (SH-083), the natural-language-to-data translator (SH-197), the report generator (SH-198) and the knowledge-base retrieval index (SH-199) — are trained and served on the NVIDIA A100/H100 HPC cluster. The cross-copilot orchestrator (SH-086) and the hash-chained logbook (SH-063) run on the in-house HPC digital-twin node, so that the fast protective loop (the clamp and the analytic monitors) never contends with the training fleet.

Grounding database. The copilots are grounded to a physics database assembled by the Kronos reference-code stack, which the predictive digital twin ^[34] mirrors: nonlinear gyrokinetic transport from CGYRO ^[37], free-boundary equilibria from a FreeGS-family solver ^[39], and neutronics/activation from OpenMC with DAGMC geometry ^[38], among the extended-MHD, edge and fast-ion codes of the wider campaign. The copilots do not re-run these solvers online; they retrieve and cite their frozen outputs, which is what keeps every answer both fast and grounded ^[23,36,35].

Problem scale and reproducibility. The clamp demonstration is a **100** × **100** trajectory-step ensemble with a closed-form per-step QP solve; the copilot training is a multi-GPU fine-tuning and index-build workload on the A100/H100 cluster. Every reported number is regenerable from the archived control-analysis script and the frozen register anchors at the stated seed. No compute-cost figures are reported here; the campaign is described by the codes it ran, not by its budget. The broader AI-and-quantum-computing methodology for the Kronos series — machine- learning surrogates, real-time control and a resource-honest quantum frontier — is set out in the companion [[doi:10.5281/zenodo.22645702](https://doi.org/10.5281/zenodo.22645702); <https://www.kronosfusionenergy.com/publications/ai-quantum-control.html>].

Results and verification

THE DEMONSTRATED CLAMP

The single demonstrated component is the clamp (table 1, figures 1–4): zero violations in **100** trajectories, peak $\beta_N = \mathbf{3.500}$ at the bound, $h_{\min} = +\mathbf{0.000}$, and a control multiplier confined to $[\mathbf{1.061}, \mathbf{1.250}]$. This is the property Theorem 3.2 propagates to the whole layer.

ACCEPTANCE CRITERIA FOR THE HUMAN-MACHINE LAYER

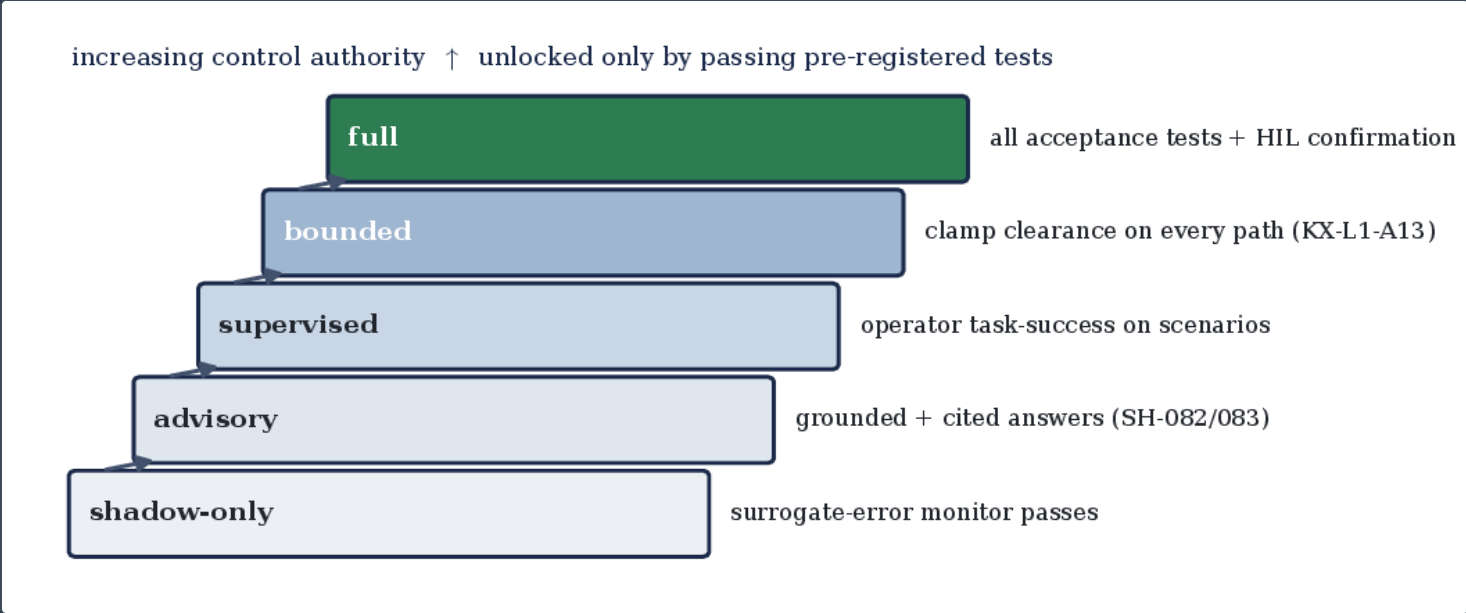
No copilot is granted operational use by assertion. Each clears an explicit, pre-registered acceptance test (table 2); the tests are the contract that gates the interface, and the clamp criterion — “no execution path bypasses the clamp” — is the non-negotiable one, guaranteed by Theorem 3.2.

@P3.6CM P3.2CM X@ COMPONENT	FUNCTION	ACCEPTANCE TEST
certified clamp (KX-L1-A13)	CBF safety filter	demonstrated: 0/100 violations, $h_{\min} = +0.000$
plasma copilot (SH-082)	NL-supervised control	grounded/cited; high intent accuracy; never bypasses clamp
engineering copilot (SH-083)	autonomous root-cause analysis	identifies seeded root causes correctly
operations copilot	shift-operations assistant	grounded/cited; operator task-success on scenarios
orchestrator (SH-086)	cross-copilot handoff	correct copilot handoff on scripted scenarios
NL→data / reports / RAG (SH-197/198/199)	grounded retrieval and reporting	query accuracy; reports trace to signed ledger
digital logbook (SH-063)	auditable system of record	every intent/answer/clamp action hash-chained

The human-machine layer components and their pre-registered acceptance criteria (register serials, ref. ^[40]). The clamp is demonstrated; the copilot and logbook criteria are specified acceptance tests built to the demonstrated pattern (§8).

MATURITY-GATED AUTHORITY

Control authority is not binary. Each component is admitted to a graded level of authority — shadow-only, advisory, supervised, bounded, full — and rises only as it clears successive acceptance criteria (figure 10). A copilot that has passed only its grounding test may advise but not command; bounded authority requires clamp clearance on every path (guaranteed here); full authority additionally requires operator task-success on scripted scenarios and a hardware-in-the-loop confirmation. Because the mapping from criteria to authority is explicit, the route from the demonstrated clamp to a fully autonomous, operator-commanded plant is a defined, testable sequence rather than an open question — the same maturity-gated principle formalised in the companion actuation-authority result [[doi:10.5281/zenodo.22645795](https://doi.org/10.5281/zenodo.22645795)].



(Schematic.) Verification-maturity gating of control authority. Each rung admits a higher level of authority and is unlocked by a pre-registered acceptance test (table 2); the “bounded” rung is the clamp-clearance guarantee of Theorem 3.2. The acceptance tests are the real bars; the ladder is a schematic of the gating policy.

Discussion

RELATION TO LEARNED CONTROL

The layer is complementary to, not a competitor of, learned control. Degraeve *et al.* [7] and Seo *et al.* [8] demonstrated learned *policies* that command a tokamak directly from measurements; the contribution here is not another policy but the guarantee that lets *any* policy be used safely. By Theorem 3.2, a learned coil-voltage or tearing-avoidance policy can generate u_{prop} and inherit the forward-invariance bound of table 1 without being trusted with authority. Relative to a purely learned controller, the layer adds the certified clamp between the policy and the plant, the grounding and abstention contract that keeps the policy's language interface honest, and the tamper-evident record that makes the whole loop auditable.

GROUNDING AGAINST THE LANGUAGE-MODEL LITERATURE

Large language models are fluent but hallucinate [27], which is exactly the failure a control room cannot tolerate. The layer's answer is architectural: retrieval-augmented generation [23] with a hard verification gate (12), instruction-tuned models [24] prompted to reason before acting [25,26], and a conformal abstention wrapper (13) [29] that hands authority back to a deterministic floor out of distribution. The transformer [21] and few-shot-capable models [22] make the natural-language surface possible; the grounding contract and the clamp are what make it safe. Crucially, a hallucinated proposal is not a safety event in this architecture: it is bounded by Theorem 3.2 and recorded by (16).

HUMAN FACTORS AND RUNTIME ASSURANCE

The human-factors literature warns that automating the routine can erode operator skill and situation awareness [17,18,19], and that the level of automation must be a deliberate design choice [20]. The layer takes those warnings as requirements: the operator supervises in natural language and confirms every actuation (the γ gate in (11)), the copilots surface grounded, cited, inspectable evidence rather than opaque commands, and the maturity ladder keeps the human's role explicit as authority is delegated. Architecturally the clamp is a control-barrier realisation of the runtime-assurance / Simplex pattern — a simple verified safety layer beneath a complex unverified one [16,12] — and the hash-chained logbook is the standard tamper-evidence construction [32]. The paper's contribution is to assemble these into a single fusion operator interface in which the safety layer is not merely present but *demonstrated*.

Limitations

One honest gate bounds the claims. The clamp is demonstrated (table 1) and the no-bypass property follows from it by Theorem 3.2; but the copilots' scripted-scenario task-success and cross-copilot handoff metrics (SH-082/083/086) are *specified* acceptance tests whose measured results are still being logged. Admitting a copilot to operational use is therefore gated behind those measurements, and the maturity ladder of figure 10 keeps each copilot advisory until its bar is met. The clamp certificate is further stated at the fidelity of a design-point-anchored β_N surrogate; the confirming step to full-plant authority is a hardware-in-the-loop bench, as for the rest of the control stack.

Conclusion

The Kronos operator interface keeps a human in command of an autonomous fusion plant by making one property demonstrated and load-bearing: copilots advise, the certified clamp decides. The control-barrier quadratic program holds β_N to its **3.5** bound with zero violations in **100** trajectories, $h_{\min} = +0.000$, and bounded control effort $u \in [1.061, 1.250]$ (table 1); the policy-agnostic no-bypass theorem (Theorem 3.2) propagates that bound to the entire layer, so the consequence of any copilot proposal is bounded whatever the copilot proposes. Above the clamp, a plasma copilot, an engineering copilot performing causal root-cause analysis, an operations copilot, a cross-copilot orchestrator, and a hash-chained digital logbook give the operator a grounded, cited, natural-language surface, with every actuation confirmed by the operator and cleared by the clamp, and every action recorded in a tamper-evident ledger. The operator stays in command by construction, and the interface is built to the demonstrated clamp-gated pattern on a staged, acceptance-gated schedule.

Acknowledgments Part of the Kronos Fusion Energy 2026 design series. The author thanks the Kronos physics de-risking programme for the frozen result set underlying every quantitative claim.

Reproducibility. The certified-clamp demonstration is a frozen anchor (KX-L1-A13) in the Kronos Physics De-Risking Register ^[40] and is regenerable from the archived control-analysis script at the stated seed (**seed = 20260726**); the copilot, orchestrator and logbook specifications and their acceptance criteria are register entries KFE-CF-SH-082, SH-083, SH-086, SH-063 and SH-197/198/199. This paper is archived at its DOI [doi:10.5281/zenodo.22645909](https://doi.org/10.5281/zenodo.22645909) and its official page <https://www.kronosfusionenergy.com/publications/the-human-machine-layer-for-autonomous-fusion-operation.html>, and the register at [doi:10.5281/zenodo.22133057](https://doi.org/10.5281/zenodo.22133057). The figure-generation script is deposited with this paper.

Data availability The clamp-demonstration summary statistics and the interface component specifications used for the figures and tables are archived with the deposit at [doi:10.5281/zenodo.22645909](https://doi.org/10.5281/zenodo.22645909), which references the Kronos 2026 series including the AI and quantum-computing paper [[doi:10.5281/zenodo.22645702](https://doi.org/10.5281/zenodo.22645702)]. All quantitative values trace to the register ^[40] by the serial identifiers cited inline. Any economic analysis is reported separately and is not part of the public deposit.

Competing interests Provisional patent applications relating to aspects of this work have been filed.

References

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 D Humphreys *et al*, Novel aspects of plasma control in ITER, *Phys. Plasmas* **22** 021806 (2015). [doi:10.1063/1.4907901](https://doi.org/10.1063/1.4907901)
- 2 J E Menard *et al*, Fusion nuclear science facilities and pilot plants based on the spherical tokamak, *Nucl. Fusion* **56** 106023 (2016). [doi:10.1088/0029-5515/56/10/106023](https://doi.org/10.1088/0029-5515/56/10/106023)
- 3 P C de Vries *et al*, Survey of disruption causes at JET, *Nucl. Fusion* **51** 053018 (2011). [doi:10.1088/0029-5515/51/5/053018](https://doi.org/10.1088/0029-5515/51/5/053018)
- 4 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** 1611 (1985). [doi:10.1088/0029-5515/25/11/007](https://doi.org/10.1088/0029-5515/25/11/007)
- 5 A Pironti and M Walker, Fusion, tokamaks, and plasma control: an introduction and tutorial, *IEEE Control Syst. Mag.* **25**(5) 30 (2005). [doi:10.1109/MCS.2005.1512794](https://doi.org/10.1109/MCS.2005.1512794)
- 6 F Felici *et al*, Real-time physics-model-based simulation of the current density profile in tokamak plasmas, *Nucl. Fusion* **51** 083052 (2011). [doi:10.1088/0029-5515/51/8/083052](https://doi.org/10.1088/0029-5515/51/8/083052)
- 7 J Degraeve *et al*, Magnetic control of tokamak plasmas through deep reinforcement learning, *Nature* **602** 414 (2022). [doi:10.1038/s41586-021-04301-9](https://doi.org/10.1038/s41586-021-04301-9)
- 8 J Seo *et al*, Avoiding fusion plasma tearing instability with deep reinforcement learning, *Nature* **626** 746 (2024). [doi:10.1038/s41586-024-07024-9](https://doi.org/10.1038/s41586-024-07024-9)
- 9 J Kates-Harbeck, A Svyatkovskiy and W Tang, Predicting disruptive instabilities in controlled fusion plasmas through deep learning, *Nature* **568** 526 (2019). [doi:10.1038/s41586-019-1116-4](https://doi.org/10.1038/s41586-019-1116-4)
- 10 C Rea *et al*, A real-time machine learning-based disruption predictor in DIII-D, *Nucl. Fusion* **59** 096016 (2019). [doi:10.1088/1741-4326/ab28bf](https://doi.org/10.1088/1741-4326/ab28bf)
- 11 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 (2017). [doi:10.1109/TAC.2016.2638961](https://doi.org/10.1109/TAC.2016.2638961)
- 12 A D Ames *et al*, Control barrier functions: theory and applications, *18th European Control Conference (ECC)* 3420 (2019). [doi:10.23919/ECC.2019.8796030](https://doi.org/10.23919/ECC.2019.8796030)
- 13 F Blanchini, Set invariance in control, *Automatica* **35**(11) 1747 (1999). [doi:10.1016/S0005-1098\(99\)00113-2](https://doi.org/10.1016/S0005-1098(99)00113-2)
- 14 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 (2000). [doi:10.1016/S0005-1098\(99\)00214-9](https://doi.org/10.1016/S0005-1098(99)00214-9)
- 15 B Stellato, G Banjac, P Goulart, A Bemporad and S Boyd, OSQP: an operator splitting solver for quadratic programs, *Math. Program. Comput.* **12**(4) 637 (2020). [doi:10.1007/s12532-020-00179-2](https://doi.org/10.1007/s12532-020-00179-2)
- 16 L Sha, Using simplicity to control complexity, *IEEE Software* **18**(4) 20 (2001).
- 17 L Bainbridge, Ironies of automation, *Automatica* **19**(6) 775 (1983). [doi:10.1016/0005-1098\(83\)90046-8](https://doi.org/10.1016/0005-1098(83)90046-8)
- 18 R Parasuraman and V Riley, Humans and automation: use, misuse, disuse, abuse, *Hum. Factors* **39**(2) 230 (1997).
- 19 M R Endsley, Toward a theory of situation awareness in dynamic systems, *Hum. Factors* **37**(1) 32 (1995).
- 20 R Parasuraman, T B Sheridan and C D Wickens, A model for types and levels of human interaction with automation, *IEEE Trans. Syst. Man Cybern. A* **30**(3) 286 (2000).
- 21 A Vaswani *et al*, Attention is all you need, *Adv. Neural Inf. Process. Syst. (NeurIPS)* **30** (2017). arXiv:1706.03762
- 22 T B Brown *et al*, Language models are few-shot learners, *Adv. Neural Inf. Process. Syst. (NeurIPS)* **33** 1877 (2020). arXiv:2005.14165
- 23 P Lewis *et al*, Retrieval-augmented generation for knowledge-intensive NLP tasks, *Adv. Neural Inf. Process. Syst. (NeurIPS)* **33** 9459 (2020). arXiv:2005.11401
- 24 L Ouyang *et al*, Training language models to follow instructions with human feedback, *Adv. Neural Inf. Process. Syst. (NeurIPS)* **35** 27730 (2022). arXiv:2203.02155

- 25 J Wei *et al*, Chain-of-thought prompting elicits reasoning in large language models, *Adv. Neural Inf. Process. Syst. (NeurIPS)* **35** 24824 (2022). [arXiv:2201.11903](https://arxiv.org/abs/2201.11903)
- 26 S Yao *et al*, ReAct: synergizing reasoning and acting in language models, *Int. Conf. on Learning Representations (ICLR)* (2023). [arXiv:2210.03629](https://arxiv.org/abs/2210.03629)
- 27 Z Ji *et al*, Survey of hallucination in natural language generation, *ACM Comput. Surv.* **55**(12) 1 (2023). [doi:10.1145/3571730](https://doi.org/10.1145/3571730)
- 28 C Guo, G Pleiss, Y Sun and K Q Weinberger, On calibration of modern neural networks, *Int. Conf. on Machine Learning (ICML)* 1321 (2017). [arXiv:1706.04599](https://arxiv.org/abs/1706.04599)
- 29 A N Angelopoulos and S Bates, Conformal prediction: a gentle introduction, *Found. Trends Mach. Learn.* **16**(4) 494 (2023).
- 30 S M Lundberg and S-I Lee, A unified approach to interpreting model predictions, *Adv. Neural Inf. Process. Syst. (NeurIPS)* **30** (2017). [arXiv:1705.07874](https://arxiv.org/abs/1705.07874)
- 31 J Pearl, Causal diagrams for empirical research, *Biometrika* **82**(4) 669 (1995). [doi:10.1093/biomet/82.4.669](https://doi.org/10.1093/biomet/82.4.669)
- 32 S Haber and W S Stornetta, How to time-stamp a digital document, *J. Cryptol.* **3** 99 (1991). [doi:10.1007/BF00196791](https://doi.org/10.1007/BF00196791)
- 33 R C Merkle, A digital signature based on a conventional encryption function, in *Advances in Cryptology — CRYPTO '87*, Lecture Notes in Computer Science **293** 369 (1988). [doi:10.1007/3-540-48184-2_32](https://doi.org/10.1007/3-540-48184-2_32)
- 34 M G Kapteyn, J V R Pretorius and K E Willcox, A probabilistic graphical model foundation for enabling predictive digital twins at scale, *Nat. Comput. Sci.* **1** 337 (2021). [doi:10.1038/s43588-021-00069-0](https://doi.org/10.1038/s43588-021-00069-0)
- 35 L Lu, P Jin, G Pang, Z Zhang and G E Karniadakis, Learning nonlinear operators via DeepONet, *Nat. Mach. Intell.* **3** 218 (2021). [doi:10.1038/s42256-021-00302-5](https://doi.org/10.1038/s42256-021-00302-5)
- 36 M Raissi, P Perdikaris and G E Karniadakis, Physics-informed neural networks, *J. Comput. Phys.* **378** 686 (2019). [doi:10.1016/j.jcp.2018.10.045](https://doi.org/10.1016/j.jcp.2018.10.045)
- 37 J Candy, E A Belli and R V Bravenec, A high-accuracy Eulerian gyrokinetic solver for collisional plasmas, *J. Comput. Phys.* **324** 73 (2016). [doi:10.1016/j.jcp.2016.07.039](https://doi.org/10.1016/j.jcp.2016.07.039)
- 38 P K Romano *et al*, OpenMC: a state-of-the-art Monte Carlo code for research and development, *Ann. Nucl. Energy* **82** 90 (2015). [doi:10.1016/j.anucene.2014.07.048](https://doi.org/10.1016/j.anucene.2014.07.048)
- 39 N C Amorisco *et al*, FreeGSNKE: a Python-based dynamic free-boundary toroidal plasma equilibrium solver, *Phys. Plasmas* **31** 042517 (2024). [doi:10.1063/5.0188467](https://doi.org/10.1063/5.0188467)
- 40 P I Ford, *Kronos Physics De-Risking Register*, Kronos Fusion Energy 2026 series. [doi:10.5281/zenodo.22133057](https://doi.org/10.5281/zenodo.22133057)

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

The Human-Machine Layer for Autonomous Fusion Operation: Natural-Language-Supervised Control, Autonomous Root-Cause Analysis and a Digital Operator Logbook

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

[Zenodo · doi:10.5281/zenodo.22645909](https://zenodo.org/doi/10.5281/zenodo.22645909)

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

[The Physics De-Risking Register](#)