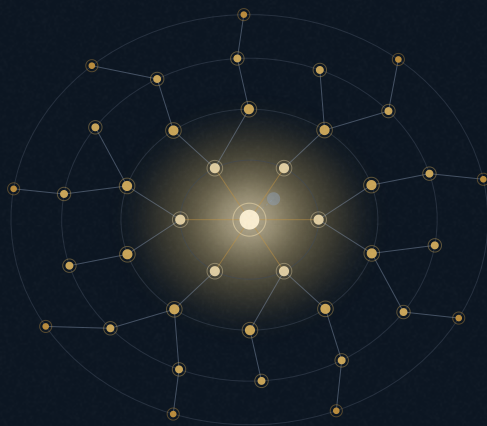




KRONOS FUSION ENERGY



PAPER 1.5 · THE KRONOS 2026 SERIES

AI and Quantum Computing for Fusion: ML Surrogates, Real-Time Control, and a Resource-Honest Quantum Frontier

A nuclear-regulated fusion generator must be controlled across ten orders of magnitude in time — microsecond magnet protection through month-scale fuel logistics — and every...

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 Several quantities in the underlying corpus moved after first publication; where a value has been withdrawn or restated, the record says so plainly.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

AI and Quantum Computing for Fusion: ML Surrogates, Real-Time Control, and a Resource-Honest Quantum Frontier

A nuclear-regulated fusion generator must be controlled across ten orders of magnitude in time — microsecond magnet protection through month-scale fuel logistics — and every automated actuation must be bounded, auditable, and reproducible. We present the umbrella account of the Kronos computational and control programme, developed to *review-article* depth, and we hold its three technology-readiness tiers strictly apart. *Physics-informed machine learning is deployable now, in the control loop, when its output is deterministically clamped*: a model-predictive controller acts inside a certified control-barrier-function (CBF) safety projection whose forward invariance we prove analytically and realise as a small quadratic program with a solver-free closed form. Over **100** Monte-Carlo trajectories the clamp admits **0** safe-set violations against **50** unclamped, holding the peak normalized pressure on the limit ($\beta_N \leq 3.5$, $h_{\min} = +0.000$, $u \in [1.06, 1.25]$); a dedicated companion study extends the verification to **5000** trajectories over three simultaneous barriers with **0** escapes. A physics-informed neural network reconstructs the negative-triangularity Grad-Shafranov equilibrium to **0.139%** RMS ψ against a free-boundary solver (**2877** $\times$ faster, **0.87 ms** inference, Grad-Shafranov residual **0.32%**); a differentiable flux operator reproduces the integrated-transport gain to $R^2 = 0.998$ (RMSE **0.088** in Q across $Q \in [1.19, 12.44]$) and recovers the frozen anchor $Q = 3.076$ as **2.998**; a graph network fuses degraded diagnostics (AUC **0.980**); and a fused strain-rate+magnetization quench-precursor detector reaches an area under the receiver-operating characteristic of **0.992**. The models are made predictive by a four-twin, **50–100 ms** shadow, carry calibrated uncertainty, and are admitted to control authority only through pre-registered acceptance inequalities. *Quantum sensing is a near-term diagnostics enhancement; quantum computation is, for fusion, an offline calibration tier with no resource cross-over this decade*: the linear kinetic (Vlasov) operator needs **11–19** logical qubits (early fault-tolerant, $\sim 2029 \pm 2$), the **6**-D kinetic and exact-alloy problems need ~ 44 and ~ 238 logical qubits, and we report a measured quantum-machine-learning *negative* (classifier AUC falls **0.817** $\rightarrow$ **0.691** as the feature map widens **4** $\rightarrow$ **8** qubits) together with a QAOA optimisation negative. The entire campaign runs on the NVIDIA GPU HPC campaign against frozen reference codes. The contribution is a certifiable, resource-honest *system and method*, not new plasma physics; every headline quantity is a frozen anchor in the de-risking register and crosses the Kronos 2026 series. Canonical anchors are the Hyperion breeder ($Q = 3.076$, $P_{\text{fus}} = 85.04 \text{ MW}$, $I_p = 9.66 \text{ MA}$, $\delta = -0.30$, centrepost **16.84 T**) and the D-³He tandem-mirror burner ($Q_E = 1.318$, plug **26.49 T**) — two distinct magnets.

Open-access deposit (code & data, CC BY 4.0): DOI [10.5281/zenodo.22645702](https://doi.org/10.5281/zenodo.22645702) · Zenodo community *kronos_fusion_energy*. Explore it live: [interactive 3-D model](#) · [physics validator](#).

Keywords: digital twin physics-informed machine learning neural operators control barrier functions model-predictive control safe reinforcement learning quantum computing quantum sensing disruption prediction uncertainty quantification

CONTENTS

Table of Contents

AI and Quantum Computing for Fusion: ML Surrogates, Real-Time Control, and a Resource-Honest Quantum Frontier	3
Introduction	7
The autonomous control architecture	12
The reduced operating-envelope plant model	18
Governing equations and the safety-projected controller	21
Formal properties of the safety-projected controller	29
Learned surrogates and calibrated uncertainty	31
An end-to-end worked control cycle	42
Quantum sensing (Tier 2, near-term)	43
Computational methodology: the NVIDIA GPU HPC campaign	47
Extended verification and the de-risking methodology	51
Diagnostics and real-time state estimation	52
The magnet-protection subsystem in depth	54
Neutronics, activation, and the fuel-cycle shadow	56
The human-machine layer	57
The burner control problem	58
The resource-honest quantum frontier	59
Sensitivity and uncertainty quantification	70
Benchmarking against published experiment and theory	72
Readiness, roadmap, and the path to full autonomy	74
Results summary	76
The companion series	79
Discussion	81
Limitations	83
Conclusion	84
Nomenclature	86
The reduced operating-envelope model	87
Standing assumptions	88
Condensed MPC quadratic program	89
KKT derivation of the closed form	90
Exponential CBF for relative-degree-two barriers	91
Ensemble-Kalman algorithm	92
PINN training and the Grad–Shafranov discretisation	93
Neural-operator architectures	94
Conformal coverage guarantee	95
Fourier–Hermite encoding of linearised Vlasov–Poisson	96
Fault-tolerant resource formulas	97
Numerical schemes and the multi-rate schedule	98
ISS proof of the clamped loop	99
NV-diamond shot-noise sensitivity	100
Explicit QUBO encodings of the three breeder problems	101
Quantum-kernel concentration	102
Assurance case and regulatory mapping	103
Derivation of the operating-envelope barriers	104

Hamilton–Jacobi reachability successor	106
Latency budget derivation	107
Diagnostic-graph construction	108
Reduced-order model for the fast loop	109
Adaptive conformal calibration	110
MLOps and provenance	111
Troyon limit from the ideal-MHD energy principle	112
Particle-filter alternative to the EnKF	113
The safe reinforcement-learning objective	114
Worked control-cycle timing	115
Bayesian optimal experimental design	116
Cross-machine transfer learning	117
Worked NV multiplexing budget	118
Error budget for the Landau circuit	119
Detection metrics	120
Conformal versus Bayesian uncertainty	121
Verification of the reference codes	122
The disruption-precursor library	123
Convergence of the Monte-Carlo verification	124
Burner actuator models	125
Comparison of quantum-simulation approaches	126
The signed-ledger data structure	127
Design decisions and their rationale	128
Register-serial glossary	129
References	134
One programme, many papers	140

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)

Introduction

FUSION IS A CONTROL PROBLEM

A compact, high-power-density fusion generator is as much a control problem as a physics problem. The Hyperion breeder holds a plasma current of $I_p = 9.66\text{ MA}$ at a major radius of only $R_0 = 1.20\text{ m}$ behind a centrepost magnet at a peak on-conductor field of 16.84 T , at aspect ratio $A = 2.5$ and elongation $\kappa = 2.0$, in a negative-triangularity ($\delta = -0.30$) regime chosen for good confinement without the edge-localised-mode cycle of conventional H-mode^[24,25]. In a device this compact the characteristic times of vertical instability, thermal transients and magnet events are short; in a nuclear system every actuation carries a safety consequence^[21,22,23]. Autonomy is therefore not a convenience but a design requirement, and an autonomous controller is only as trustworthy as the model beneath it and the floor beneath that model.

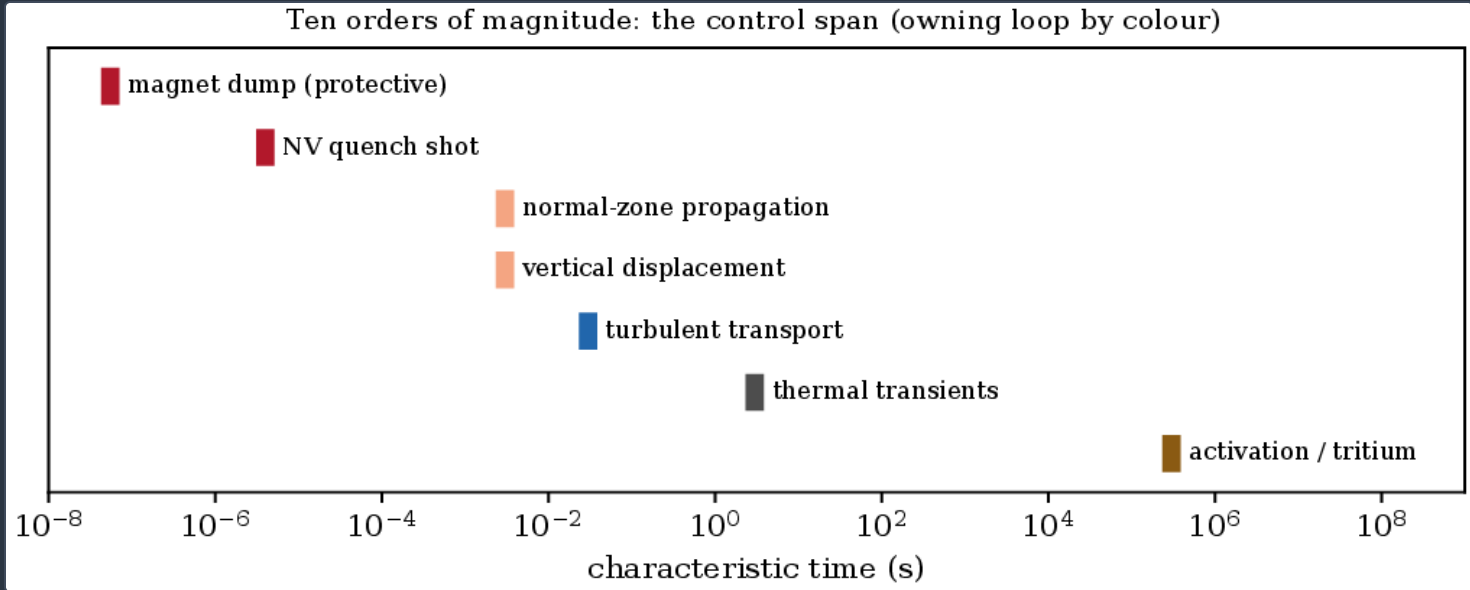
The compactness that makes the spherical tokamak an attractive breeder is exactly what raises the control bar. At aspect ratio $A = 2.5$ and major radius $R_0 = 1.20\text{ m}$ the plasma sits close to the centrepost, the field line pitch is steep, and the elongation $\kappa = 2.0$ places the plasma near the vertical-stability boundary, where a small displacement grows on the resistive-wall time unless actively controlled. High power density in a small volume means the thermal and neutron loads are concentrated, so the margins on the first wall, the divertor, and the centrepost are tighter than in a large machine and leave less time to react to an excursion. Negative triangularity buys a quiet, edge-localised-mode-free edge^[24,25], but only inside a narrow operating window the controller must actively hold; step outside it and the confinement benefit is lost. Every one of these features converts a physics advantage into a control requirement, which is why the computational and control programme is not an add-on to the machine but a co-equal part of its design.

The nuclear context sharpens every one of these requirements. A conventional research tokamak can tolerate a disruption as an expensive but survivable event; a nuclear-regulated generator cannot, because a disruption in a compact, high-field, high-power-density machine threatens the first wall, the divertor, and the centrepost magnet, and because every automated actuation carries a regulatory consequence that must be bounded, auditable, and reproducible. The controller is therefore not merely a performance optimiser but a safety-critical component, and its trustworthiness must be established with the same rigour as the mechanical integrity of the vessel. This is the sense in which the programme is verification-first: the burden of proof is on the controller to demonstrate it cannot leave the safe envelope, not on the operator to hope that it will not.

The magnitude of the timescale span is the organising difficulty. A vertical displacement event grows on the resistive-wall time, $\mathcal{O}(1\text{ ms})$; a developing quench in a no-insulation high-temperature-superconducting winding propagates on the normal-zone timescale, $\mathcal{O}(1\text{ ms} - 10\text{ ms})$; the magnet-dump protective action must complete in $\mathcal{O}(100\text{ ns} - 1\text{ }\mu\text{ s})$; turbulent transport equilibrates on $\mathcal{O}(1\text{ ms} - 100\text{ ms})$; thermal transients in the structure evolve over seconds; activation and tritium inventory bookkeeping evolve over hours to months. No single controller, and no single computer, spans that range: the architecture must be layered by timescale and by security domain, with the fastest, most safety-critical loop carried by the simplest, most deterministic hardware. Table 1 sets out the span and the loop that owns each band.

PHENOMENON	TIMESCALE	OWNING LOOP
magnet dump (protective)	0.055 μs	LO hardware failsafe
NV quench precursor shot	4 μs	fast protective
normal-zone propagation	ms	fast protective
vertical displacement event	ms	reduced-order twin
turbulent transport	ms–100 ms	neural flux operator
thermal transients	s	thermomechanics twin
activation / tritium inventory	hours–months	neutronics shadow

The ten-orders-of-magnitude control span and the loop that owns each band. The fastest, most safety-critical action is carried by the simplest, most deterministic hardware.



The ten-orders-of-magnitude control span of table 1 on a logarithmic time axis, from the $0.055\mu s$ magnet dump to month-scale inventory bookkeeping. Colour indicates the owning loop; the fastest, most safety-critical action is carried by the simplest hardware.

Feedback control of tokamak equilibria and profiles is a mature discipline — magnetic reconstruction feeds real-time equilibrium ^[29] and shape controllers ^[30], and physics-based profile simulators run inside the control cycle ^[31]. More recently, learned controllers have entered the loop: Degraeve *et al.* trained a deep reinforcement-learning (RL) agent to command TCV coil voltages directly from magnetics, including negative-triangularity shapes ^[32]; Seo *et al.* used RL to steer DIII-D clear of a tearing instability ^[33]; and deep networks forecast disruptions across machines ^[34,35]. These advances share a thread: control quality is bounded by the speed and fidelity of the underlying model, and none of them, by itself, supplies the certified last line of defence a regulated plant requires.

DESIGN PHILOSOPHY: VERIFICATION-FIRST AUTONOMY

The programme rests on three principles that recur through every section. *First, safety is a boundary, not a behaviour.* Rather than trying to train a controller that is safe, we place a deterministic, provably-invariant boundary beneath any controller, so that safety is a property of the closed loop that holds regardless of what the learned layers do (§4). *Second, authority follows evidence.* No learned component commands an actuator until it clears a pre-registered acceptance inequality; the mapping from evidence to authority is explicit and machine-checked, so the system is auditable by construction (§2, §19). *Third, every number is frozen and traceable.* Each quantitative claim originates from a reference code, is stamped as an anchor in the de-risking register with a serial identifier, and is regenerable from archived scripts — so the paper is a set of reproducible results, not assertions (§9). These principles are what make an autonomous nuclear controller a subject an auditor and an operator can reason about, and they are the reason the quantum tiers are reported with their negatives rather than their promise: an honest map is worth more to a regulated programme than an optimistic one. The three principles are not independent; they reinforce one another. Safety-as-a-boundary is only meaningful if the boundary’s assumptions are verified, which is authority-follows-evidence; and authority-follows-evidence is only auditable if every number is frozen and traceable, which is the third principle. Together they turn a collection of learned components into a system whose behaviour can be bounded, whose authority can be justified, and whose claims can be reproduced — the three things a regulator of an autonomous plant must be able to check.

THREE READINESS TIERS, HELD APART

Machine learning and quantum technologies are routinely invoked for fusion without a clear statement of what is deployable now versus what is a decade away. This paper is the umbrella of the Kronos computational programme and holds three tiers strictly apart, reporting each with its evidence and its misses:

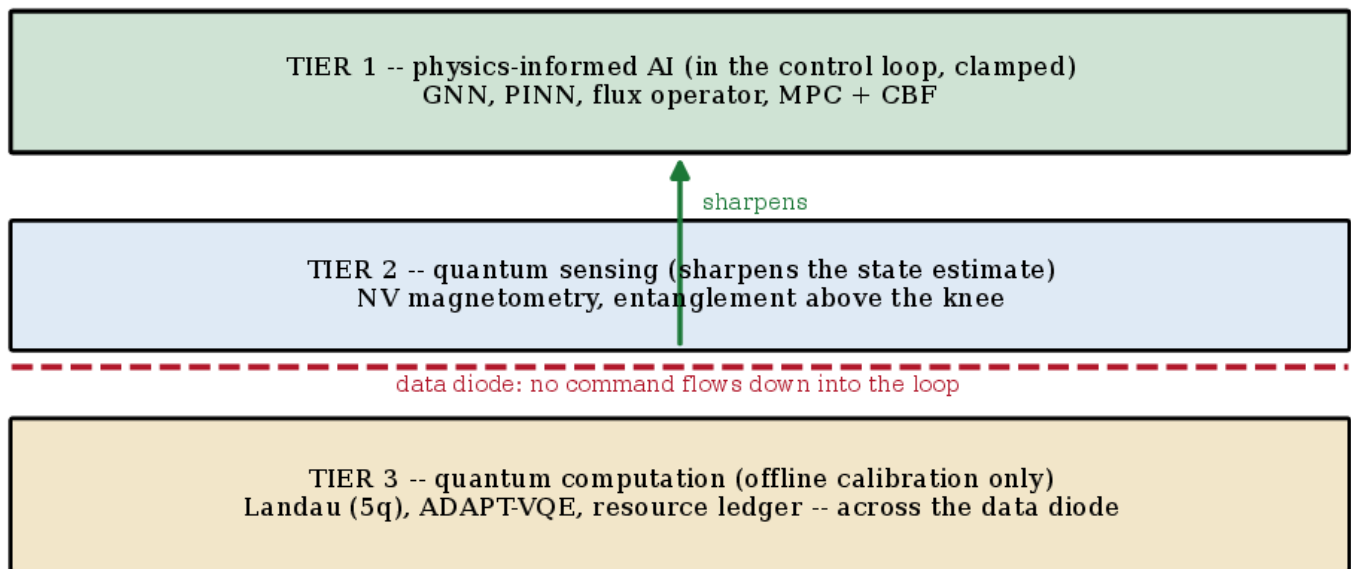
Tier 1 — physics-informed AI, deployable now if clamped. Learned surrogates and optimisation-based controllers run in the loop *because* their error is bounded and deterministically clamped by a certified safety projection. This is the load-bearing contribution.

Tier 2 — quantum sensing, near-term diagnostics. Magnetometry at or near the quantum limit sharpens the state estimate that feeds the loop; it enhances the sensing layer rather than replacing it.

Tier 3 — quantum computation, an offline calibration tier. For fusion there is no resource crossover this decade. We give dated fault-tolerant resource estimates and report measured quantum-machine-learning and quantum-optimisation *negatives* rather than reframing them. itemize

The separation is not rhetorical. Tier 1 is admitted to the control loop; Tier 2 sharpens the estimate that Tier 1 consumes; Tier 3 sits entirely offline, informing calibration and materials selection but never touching an actuator. The three are separated physically, by security domain, and by an explicit assurance argument that binds runtime authority to verification credibility. The reason to insist on the separation is that the failure modes of the three tiers are utterly different and must not be conflated. A Tier-1 surrogate can be wrong, and the certified clamp bounds the consequence; a Tier-2 sensor can degrade, and the model-fusion floor bounds the consequence; a Tier-3 quantum result can be a measured negative, and reporting it as such is the consequence. Collapsing the tiers — treating a quantum-computing promise as if it were a deployable capability, or a learned surrogate as if it were a certified controller — is exactly the error a regulated programme cannot afford, and the architecture is built to make that error structurally impossible: a Tier-3 result cannot reach an actuator because it is on the far side of the data diode, and a Tier-1 surrogate cannot command outside its authority because the maturity gate forbids it (figure 2).

(Schematic.) Three readiness tiers, held apart



(Schematic.) The three readiness tiers, held apart. Tier 1 (physics-informed AI) is in the control loop, clamped; Tier 2 (quantum sensing) sharpens the state estimate Tier 1 consumes; Tier 3 (quantum computation) sits offline across the data diode and never reaches an actuator.

PRIOR ART, IN DEPTH

Because this paper heads a series and reviews a whole computational programme, we set the prior art out carefully across the five threads it touches, so that the contribution can be located precisely against each.

Classical tokamak control.

Model-based tokamak control is mature. Real-time equilibrium reconstruction from magnetics — the EFIT lineage ^[29] — feeds shape and vertical-position controllers built on linearised plasma-response models ^[30]; physics-based profile simulators such as RAPTOR run inside the control cycle and enable current-profile control ^[31]. These controllers are trustworthy precisely because they are simple and their models are transparent, but they are bounded by the fidelity and the speed of those reduced models: a linearised response is accurate only near the operating point, and a profile solver fast enough for the loop is necessarily coarse. The Kronos architecture keeps this classical backbone — magnetic reconstruction, shape and profile controllers, a reference governor — and augments it with learned operators that extend the model's reach and speed, then places a certified boundary beneath the whole stack. This is a deliberately conservative choice: the classical backbone is retained because it is transparent and well-understood, and the learned layers are added on top of it rather than in place of it, so a failure of a learned component degrades the system toward the classical controller and ultimately to the deterministic floor, never toward an un-analysable state. The learned operators earn their place by being faster and more faithful than the reduced models the classical backbone would otherwise use, but they never become load-bearing for safety — that role is reserved for the proven boundary.

Learned control.

Deep reinforcement learning has entered the tokamak control loop. Degraeve *et al.* trained an RL agent to command TCV coil voltages directly from magnetic measurements and produced a range of shapes including negative triangularity and a droplet configuration ^[32]; Seo *et al.* trained an agent to steer DIII-D clear of a tearing instability while sustaining performance ^[33]; and disruption forecasting by deep networks now transfers across machines ^[34,35]. These are learned *policies* acting on the present measurement, and they are impressive demonstrations of performance. What none supplies, by construction, is an unconditional safety guarantee an auditor can check: an RL policy is a function whose safe behaviour is established empirically, not proven. The Kronos contribution is orthogonal and composable — a certified boundary that lets any such policy pursue performance while the machine provably cannot leave its safe set.

Digital twins and surrogates.

A digital twin is a model of a specific asset, continuously updated by that asset's data and used to predict and steer it ^[50]. Its predictive value depends on running faster than the plant. Three developments make a real-time predictive twin feasible: neural operators — DeepONet ^[44] and the Fourier neural operator ^[45], with physics-informed training ^[42,43] — that emulate an expensive solver at control latency, and that in fusion already run quasilinear transport three-to-five orders of magnitude faster than the code they replace ^[46,47,48]; projection-based model-order reduction giving certified $\sim 100\times$ reduced models for the fastest loops ^[49]; and ensemble data assimilation — the ensemble Kalman filter ^[56,57,58] — that propagates a calibrated state distribution so epistemic uncertainty travels with the estimate. Prior art treats these ingredients separately; the four-twin architecture couples them, makes each domain differentiable, carries uncertainty end-to-end, and gates each into authority by an explicit test.

Safety-critical control theory.

Control barrier functions provide forward invariance of a safe set through a pointwise inequality solved as a small QP ^[39,40], and set-invariance theory gives the underlying Nagumo condition ^[41]; model-predictive control provides constrained performance with stability guarantees ^[36,37] and fast QP solvers make it real-time ^[38]. The novelty here is not the CBF or the MPC in isolation but their composition in a nuclear loop: a solver-free closed-form CBF projection beneath a learned MPC, with an ISS argument that makes safety unconditional in the surrogate error, a robust tightening sized to a declared disturbance bound, and a decoupling into scalar projections that keeps the clamp constant-time on the hot path.

Quantum computing and sensing for plasmas.

Quantum computing is often invoked for fusion without a disciplined resource account. The literature is in fact clear: the linear kinetic (Vlasov) response is a legitimate Hamiltonian-simulation target on early-fault-tolerant machines ^[81,82,83], variational chemistry reaches the alloy fragments at chemical accuracy today ^[72,73,71,74], and the obstacles — barren plateaus ^[68], near-term limits ^[65,66,67], and the absence of demonstrated exponential advantage in ground-state chemistry ^[76] — are documented. Quantum sensing is nearer-term: NV-diamond magnetometry at the picotesla scale is a mature technology ^[60,61,62], and quantum metrology gives the Heisenberg-limit scaling that pays only in the sensitivity-limited regime ^[63]. This paper's stance is to test each claim rather than assume it, and to report the negatives.

CONTRIBUTION AND THE SERIES IT HEADS

The contribution is a *system and a method*: (i) an eight-layer, four-twin autonomous control architecture made predictive by learned neural operators and carrying calibrated uncertainty end to end (§2); (ii) a first-principles account of the reduced operating-envelope plant model and its safety barriers (§3); (iii) a governing-equation account of the safety-projected controller — the model-predictive-control (MPC) quadratic program (QP) bounded by a control-barrier-function projection with an analytic forward-invariance proof (§4); (iv) the learned-surrogate and uncertainty machinery — physics-informed and graph learning, ensemble assimilation, conformal and out-of-distribution (OOD) gating (§6); (v) quantum sensing at the quantum limit (§8); (vi) a reproducible computational methodology on the NVIDIA GPU HPC campaign (§9); (vii) a resource-honest quantum frontier with a dated roadmap (§16); and (viii) sensitivity, uncertainty-quantification and benchmarking studies (§17–§18). Every quantitative claim is a frozen anchor in the Kronos Physics De-Risking Register ^[1]; serial identifiers (e.g. KX-L1-A13) are given inline so each number is auditable, and each sub-topic has a dedicated companion paper in the 2026 series, cross-cited by reserved DOI and official page in §24. The contribution is not new plasma physics.

It is worth being precise about what is and is not claimed. We do *not* claim a new confinement concept, a new transport model, or a quantum advantage; the plasma physics is the frozen design point of the companion register, and the quantum results include measured negatives. We *do* claim a certifiable, resource-honest system and method: a control architecture in which a learned performance layer is bounded by a proven safety layer, made predictive by learned neural operators, carrying calibrated uncertainty end-to-end, gated into authority by pre-registered acceptance inequalities, and verified against frozen reference codes on the NVIDIA GPU HPC campaign — with the quantum tiers placed exactly where the evidence puts them. The novelty is in the composition and the discipline, not in any single component: control-barrier functions, model-predictive control, neural operators, ensemble assimilation, conformal prediction, NV magnetometry, and variational quantum algorithms are all established, and the contribution is to assemble them into a system a regulator and an operator can reason about, and to report the whole — including the misses — with reproducible rigour.

NOTATION AND CONVENTIONS

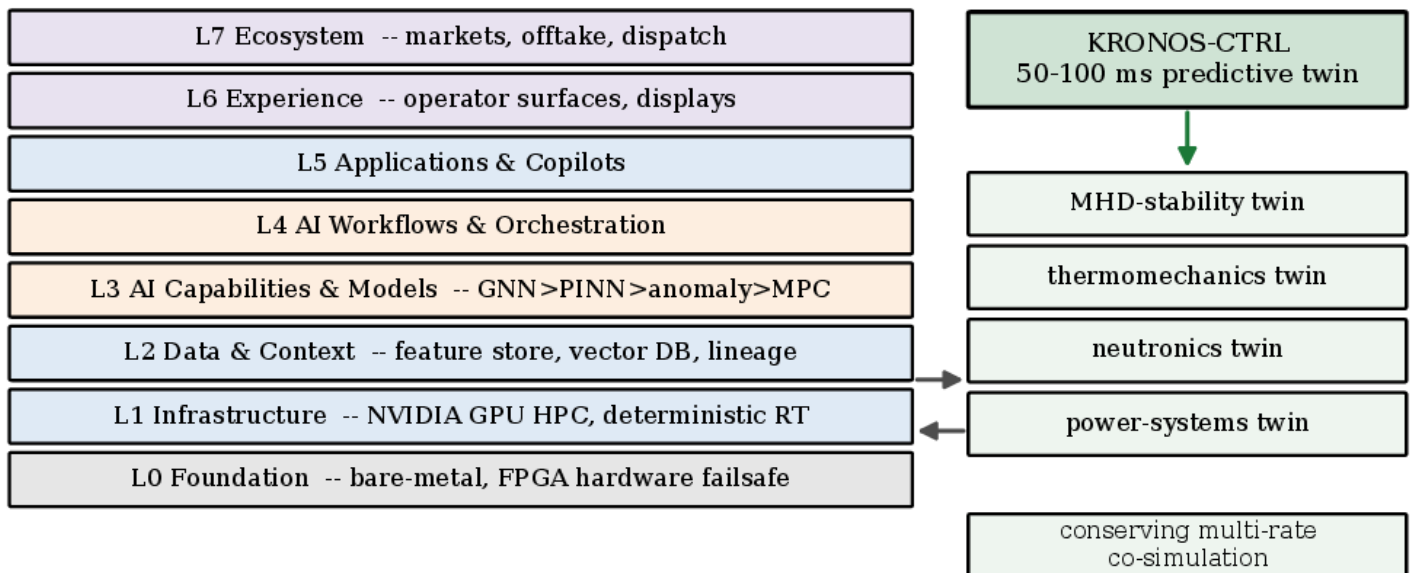
Vectors are bold lowercase ($\mathbf{x}$), matrices bold uppercase ($\mathbf{P}$); $\|\cdot\|$ is the Euclidean norm and $\|\mathbf{v}\|_Q^2 = \mathbf{v}^\top \mathbf{Q} \mathbf{v}$; $[\cdot]_+ = \max(0, \cdot)$; $L_{\mathbf{f}}\mathbf{h} = \nabla \mathbf{h} \cdot \mathbf{f}$ is the Lie derivative of $\mathbf{h}$ along $\mathbf{f}$; an extended class- $\mathcal{K}$ function $\alpha: \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0}$ is continuous, strictly increasing, and $\alpha(0) = 0$. Table 25 in Appendix 25 collects the symbols. All physics quantities are the Tier-2 frozen anchors of Table 20; where a number carries a register serial we cite it inline. No economic quantity appears anywhere in this paper by design.

The autonomous control architecture

EIGHT LAYERS AND THE LAST LINE OF DEFENCE

The stack spans eight numbered layers, from bare-metal and field-programmable-gate-array (FPGA) hardware at the foundation (L0), through infrastructure and data (L1–L2), the AI capabilities and models (L3), workflows and copilots (L4–L5), to the experience and ecosystem layers (L6–L7) that face operators and markets (figure 3). The identical architecture drives both machines: only the physical plant, the actuators, and the physics models differ between the Hyperion breeder and the D–³He tandem-mirror burner. The organising principle is that *no AI output reaches an actuator without passing a deterministic, hardware-enforced clamp*. The learned layers pursue performance; a proven boundary beneath them guarantees the machine stays in its safe envelope even when a model is wrong. This inverts the usual burden of proof for a learned controller: rather than establishing empirically that the controller behaves safely across the operating space — a claim no finite test set can fully support — the architecture establishes analytically that the boundary cannot be crossed, and lets the controller be as learned and as capable as the operator wishes above it. That boundary — deterministic rules plus a hardware fail-safe, exercised in **0.055 μ s** at the magnet dump — needs neither the network nor the cloud, so loss of the AI degrades the plant to a safe state rather than an unsafe one. The eight layers are the scaffolding that carries this principle from the actuator up to the market interface, with the security domains cutting across them so that the trust required at each layer is matched to the consequence of its failure.

(Schematic.) Eight-layer control stack and the four coupled twins



deterministic safety boundary (hardware clamp -- the last line of defense)

(Schematic.) The eight-layer control stack (L0–L7), the four coupled twins, and the KRONOS-CTRL **50–100 ms** predictive twin. A deterministic safety boundary sits below the learned layers: the hardware clamp is the last line of defence. The figure encodes structure only; no data values are shown.

The layered structure is a direct response to the timescale span of §1. Layer L0 hosts the protective loop — deterministic combinational logic and a hardware failsafe, no operating system, no dynamic memory allocation, no network dependency — so that the microsecond magnet dump is provably bounded in latency. Layers L1–L2 host the data plane: a feature store, a time-series historian, a vector database for retrieval, and a lineage service that timestamps every training example and every runtime decision. Layer L3 is where the learned models live — the graph-network reconstructor, the physics-informed equilibrium and flux operators, the anomaly-ensemble, and the MPC optimiser. Layers L4–L5 orchestrate workflows and expose operator copilots. Layers L6–L7 are the operator experience and the market/dispatch interface. The security domains cut across the layers: the protective loop is air-gapped from the training fleet, joined only by a physical data diode, so a compromised or overloaded training job can never delay a protective action. Table 2 states each layer’s role, its dominant timescale, and its security domain.

LAYER	ROLE	TIMESCALE	DOMAIN	
L7 ecosystem	markets, offtake, dispatch	hours–days	enterprise	
L6 experience	operator surfaces, displays	seconds	operations	
L5 applications \	copilots	workflows, natural-language supervision	seconds	operations
L4 AI workflows		orchestration, scheduling	sub-second	operations
L3 AI models	GNN → PINN → anomaly → MPC	milliseconds	control	
L2 data \	context	feature store, vector DB, lineage	milliseconds	control
L1 infrastructure	NVIDIA GPU HPC, deterministic RT nodes	μs–ms	control	
L0 foundation	bare-metal, FPGA hardware failsafe	0.055 μs	protective (diode)	

The eight-layer control stack: role, dominant timescale, and security domain. The protective loop (L0) is air-gapped from the training fleet and joined only by a physical diode.

THE FOUR COUPLED TWINS AND THE PREDICTIVE SHADOW

The plant is shadowed by four digital twins, each owning a physical domain and exposing a small set of interface quantities (table 3): an MHD-stability twin (free-boundary equilibrium and ideal-MHD margins — vertical stability at $\kappa = 2.0$, the external-kink and ballooning distances, the resistive-wall response), a thermomechanics twin (the heat path and the structure — first-wall and divertor surface loading, the blanket temperature field, the centrepost’s structural and thermal margins), a neutronics twin (the 14 MeV source and everything downstream — tritium breeding, activation inventory and shutdown dose ^[55]), and a power-systems twin (magnets, current leads, cryogenics, and the electrical and quench state of the coil set; its magnet subsystem is the demonstrated exemplar of §12). The twins are coupled through a conserving multi-rate co-simulation and made predictive by learned flux and equilibrium neural operators, so the controller acts on where the plasma *will* be, not where it was; the shadow lead-time is held at $\tau_{\text{lead}} \geq 50\text{ ms}$. The full formalism is the subject of the companion four-twins paper ^[2]; here we state the pieces the umbrella needs and give the equations.

TWIN	PHYSICAL DOMAIN	ACCEPTANCE METRIC
MHD-stability	equilibrium, kink/ballooning, VS	RMS $\psi < 0.5\%$ vs free-boundary (SH-031)
thermomechanics	surface/blanket heat, centrepost	margins within fabrication-FEA fidelity
neutronics	source, TBR, activation, dose	tracks the operating point in shadow
power-systems	magnets, cryo, quench state	magnet twin demonstrated (§12)

The four twins, their domains, and the real-time acceptance metric (serials in the register ^[1]).

Coupling boundary quantities.

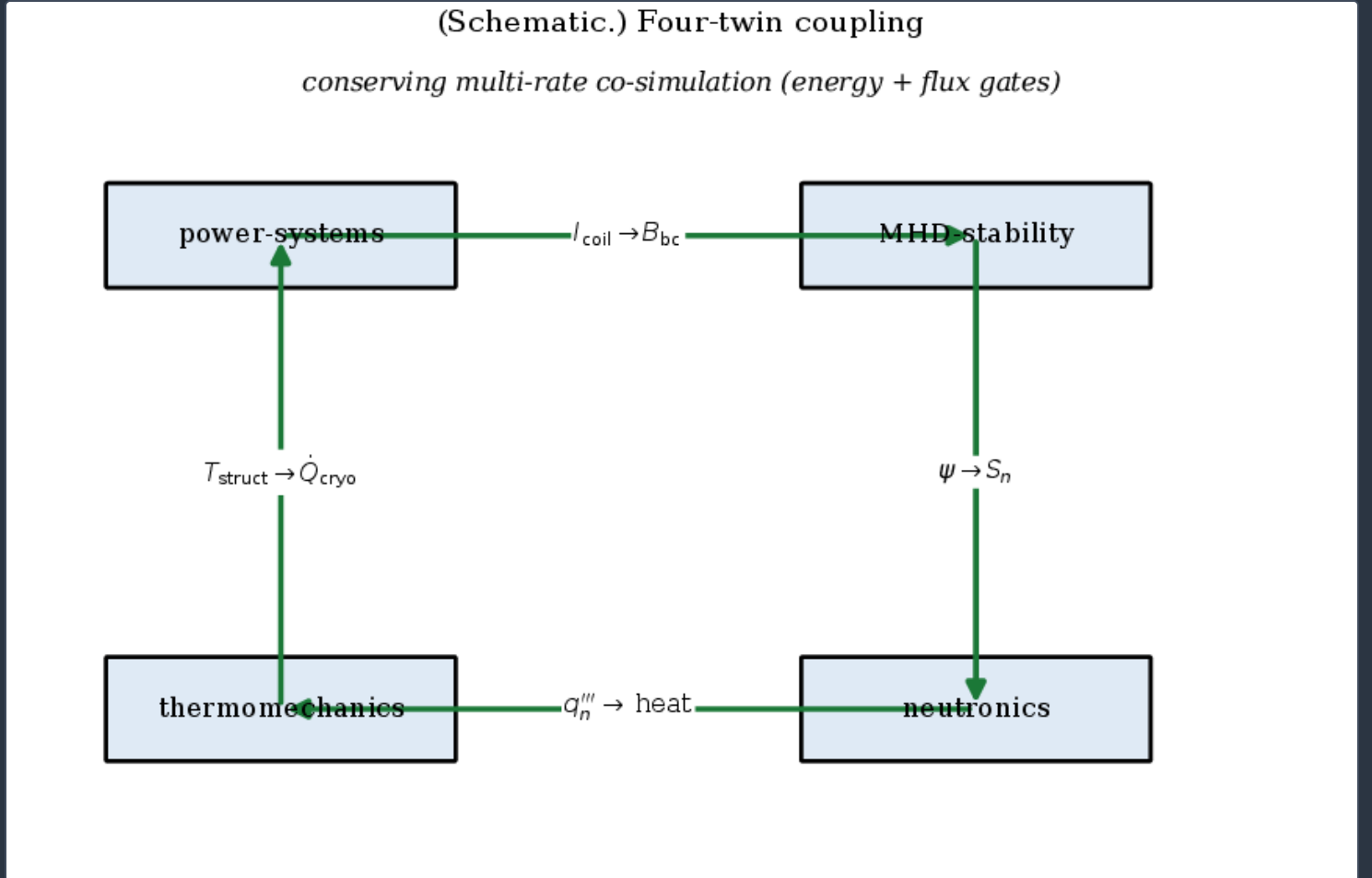
Writing $\mathbf{x} = (\mathbf{x}_M, \mathbf{x}_T, \mathbf{x}_N, \mathbf{x}_P)$ for the sub-states of the MHD, thermomechanics, neutronics and power-systems twins, the dominant couplings are the directed edges

$$\begin{aligned} \text{power} \rightarrow \text{MHD:} & \quad g_{P \rightarrow M} = I_{\text{coil}} \mapsto \text{boundary field } B_{\text{bc}}, \\ \text{MHD} \rightarrow \text{neutronics:} & \quad g_{M \rightarrow N} = \psi(R, Z) \mapsto \text{source geometry } S_n, \\ \text{neutronics} \rightarrow \text{thermo:} & \quad g_{N \rightarrow T} = q_n''' \mapsto \text{volumetric nuclear heating}, \\ \text{thermo} \rightarrow \text{power:} & \quad g_{T \rightarrow P} = T_{\text{struct}} \mapsto \text{cryogenic load } \dot{Q}_{\text{cryo}}, \end{aligned}$$

so each twin evolves under its own dynamics driven by actuators $\mathbf{u}_i$ and by coupling inputs $\mathbf{c}_i = \{g_{j \rightarrow i}\}$ from its neighbours,

$$\frac{d}{dt} \mathbf{x}_i = \mathbf{f}_i(\mathbf{x}_i, \mathbf{u}_i, \mathbf{c}_i), \quad i \in \{M, T, N, P\}.$$

The coupling edges 1–1 form the directed graph of figure 4.



(Schematic.) The four-twin coupling graph of 1–1: coil currents set the boundary field for the MHD twin, the equilibrium sets the neutron source geometry, the nuclear heating drives the thermomechanics twin, and the structural temperature sets the cryogenic load, closing the loop through the conserving multi-rate co-simulation.

Multi-rate co-simulation and conservation gates.

The domains have widely separated natural timescales, so the orchestrator advances twin i at its own step h_i and reconciles on a common macro-step $\Delta t_{\text{macro}} = n_i h_i$. Multi-rate integration is admissible only if the coupling neither creates nor destroys the shared invariants between reconciliation points; we impose two conservation gates on every macro-step $[t, t + \Delta t_{\text{macro}}]$. With E_i the energy stored in domain i and $W_{\text{ext}}, Q_{\text{ext}}$ the external work and heat crossing the plant boundary,

$$\left| \sum_i [E_i(t + \Delta t_{\text{macro}}) - E_i(t)] - (W_{\text{ext}} + Q_{\text{ext}}) \right| \leq \varepsilon_E E_{\text{ref}},$$

and at the power↔MHD interface the poloidal-flux linkage must be consistent,

$$\left\| \psi_M[I_{\text{coil}}]_{\partial\Omega} - \psi_{\text{bc}} \right\|_2 \leq \varepsilon_\Phi \|\psi\|_2.$$

A macro-step that violates either gate is rejected and retaken at a smaller Δt_{macro} . Equations 3–4 are the machine-checkable form of the acceptance criterion “energy/flux conserved across couplings” (SH-096). Appendix 37 details the schedule and the stability condition on the macro-step.

Differentiable operators and on-the-fly linearisation.

A twin is predictive only if it is faster than the plant while remaining faithful, and controllable only if it exposes its sensitivities. Both learned operators are $\mathcal{C}^1$ and are differentiated by reverse-mode automatic differentiation, so about the current shadow state $\mathbf{x}_k$ the twin linearises as

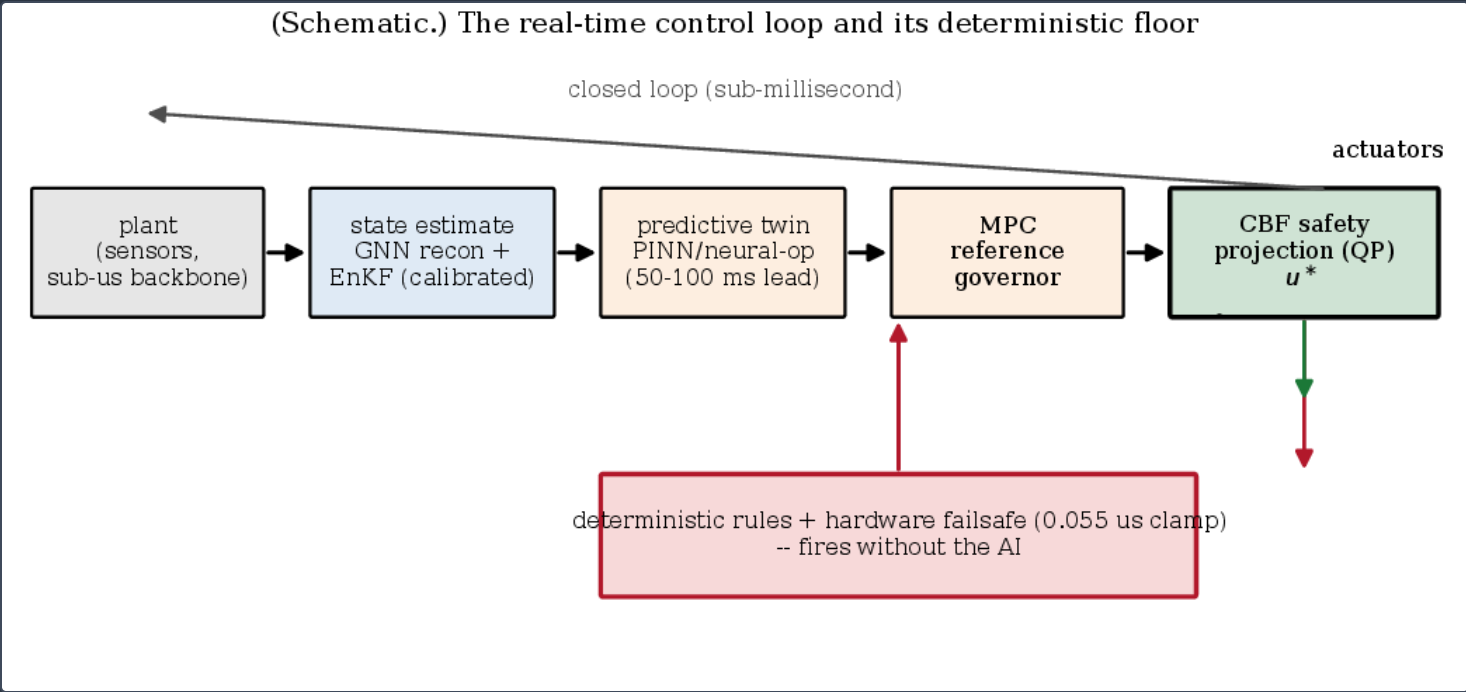
$$\mathbf{x}_{k+1} \approx \mathcal{S}_\theta(\mathbf{x}_k, \mathbf{u}_k) + \mathbf{J}_k(\mathbf{x} - \mathbf{x}_k), \quad \mathbf{J}_k = \frac{\partial \mathcal{S}_\theta}{\partial \mathbf{x}} \bigg|_{\mathbf{x}_k},$$

supplying the local linear model the controller uses in 14 without a finite-difference sweep. The differentiable twin is accepted only when its AD Jacobian matches a finite-difference reference, $\|\mathbf{J}_k^{\text{AD}} - \mathbf{J}_k^{\text{FD}}\|_\infty < 10^{-3}$ (SH-032). The shadow lead-time is $\tau_{\text{lead}} = \tau_{\text{advance}} - \tau_{\text{pipeline}}$: with the reduced-order twin advancing at $\geq 100\times$ real time (SH-093) and the medium-loop pipeline at $\mathcal{O}(5\text{ ms})$ (Appendix 46), $\tau_{\text{lead}} \geq 50\text{ ms}$ is met with margin, so the controller acts on where the plasma will be, not where it was.

THE MODEL PIPELINE: GNN → PINN → ANOMALY → MPC

The real-time loop (figure 5) chains four learned or optimisation-based stages inside a sub-millisecond budget. A graph neural network (GNN) reconstructs the plant state from the diagnostic set and imputes dropped channels; a physics-informed neural network (PINN) supplies a fast, differentiable equilibrium and the neural-operator transport surrogate supplies the fluxes; an anomaly-ensemble flags precursors; and an MPC reference governor computes the nominal actuation. Every candidate action then passes the CBF safety projection before it reaches an actuator. The estimate that enters the loop is calibrated: a state distribution, not a point, so the safety layer can read how far to trust the model.

The ordering of the pipeline is deliberate and each stage feeds the next a strictly more actionable object. The GNN turns raw, partially-degraded diagnostics into a complete state estimate with reconstructed channels; the PINN and flux operator turn that state into a fast, differentiable forecast of where the plasma will be; the anomaly-ensemble turns the forecast into an off-normal score and a calibrated uncertainty; the MPC turns the forecast and the target into a nominal command; and the CBF projection turns the nominal command into a certified-safe command. At no point does an uncalibrated or unverified object reach an actuator: the uncertainty is carried end-to-end (the GNN and ensemble emit distributions, the conformal wrapper calibrates them, the gate reads them), and the safety projection is the final, deterministic transformation. The whole pipeline fits inside a sub-millisecond budget because each stage is either a single forward pass (GNN, PINN, flux) or a closed-form/warm-started solve (MPC condensed QP, CBF projection), and the protective loop beneath it runs asynchronously in microseconds (§9).



(Schematic.) The sub-millisecond real-time control loop: calibrated state estimate → predictive twin → MPC reference governor → CBF safety projection → actuator, with a deterministic rules-plus-hardware floor that fires without the AI. Structure only; no data plotted.

MATURITY-GATED ACTUATION AUTHORITY

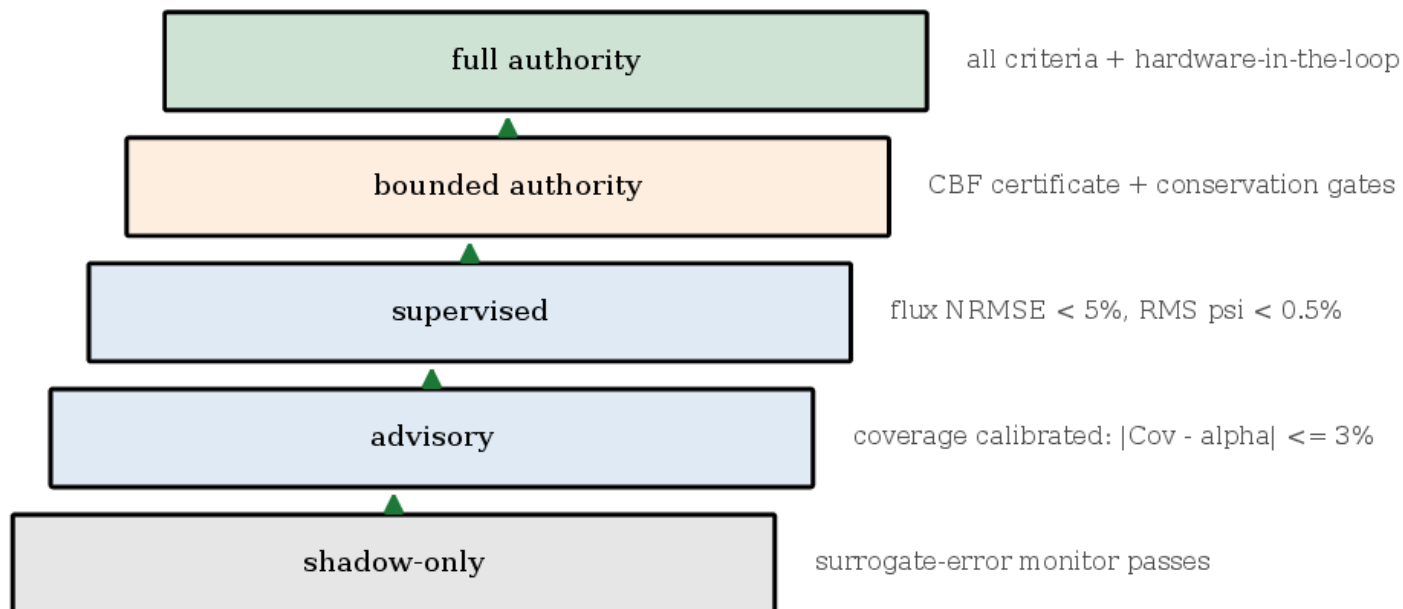
Control authority is not binary. Each learned component is admitted to a graded level of authority — shadow-only, advisory, supervised, bounded, full — and rises only as it clears successive, pre-registered acceptance inequalities (figure 6); the mapping from evidence to authority is explicit, so a model that has passed only its surrogate-error monitor may inform but not command ^[11]. This binds runtime authority to verification credibility and makes the path from a demonstrated subsystem to a fully autonomous plant a defined, testable sequence rather than an open question. The gate is the runtime expression of the assurance case (§24). The five rungs are explicit (table 4): *shadow-only* (the component runs but its output is logged, never acted on), *advisory* (its output is shown to the operator but not the controller), *supervised* (it informs the controller under operator oversight), *bounded authority* (it commands within the CBF-certified envelope), and *full authority* (it commands under all acceptance criteria plus hardware-in-the-loop confirmation). A component rises one rung only by clearing the pre-registered inequality that unlocks it, and it is demoted automatically when a runtime monitor — the OOD gate, the conformal coverage check, or the predictive-variance threshold — fires.

RUNG	WHAT THE COMPONENT MAY DO	UNLOCKING EVIDENCE
shadow-only	run and log, never act	surrogate-error monitor passes
advisory	inform the operator	conformal coverage calibrated (SH-027/073)
supervised	inform the controller, overseen	flux NRMSE < 5%, RMS ψ < 0.5% (SH-026/031)
bounded authority	command within the safe set	CBF certificate + conservation gates
full authority	command unrestricted	all criteria + hardware-in-the-loop

The five authority rungs and the evidence that unlocks each. Authority rises only by clearing the pre-registered bar and is demoted automatically when a runtime monitor fires.

(Schematic.) Maturity-gated actuation authority

control authority rises only as acceptance inequalities are cleared



(Schematic.) Maturity-gated authority. Each rung admits a higher level of control authority and is unlocked by the pre-registered acceptance inequalities of §6. The criteria are the real, frozen bars; the authority ladder is a schematic of the gating policy.

The reduced operating-envelope plant model

The certified safety layer of §4 acts on a reduced, control-affine model of the operating point. We derive that model here and construct its three safety barriers from first principles, because the guarantee is only as meaningful as the model it is stated over. The full gyrokinetic and MHD physics live in the frozen reference codes (§9); the reduced model is a conservative operating-envelope proxy against which the clamp is proved, and whose caps are declared below the hard physical limits.

STATE AND ACTUATION

Let the reduced operating state be

$$\mathbf{x} = (f_G, \beta_N, q_{95})^\top \in \mathbb{R}^3,$$

the Greenwald density fraction f_G , the normalized pressure β_N , and the edge safety factor q_{95} . These three coordinates carry the three operational limits that a compact ST must respect: the density limit, the pressure (β) limit, and the current/kink limit. The actuation vector $\mathbf{u}$ collects the fuelling rate, the auxiliary heating power, and the plasma-current/coil-current command, mapped to $(\dot{f}_G, \dot{\beta}_N, \dot{q}_{95})$ through the actuator model of Appendix 26. The reduced dynamics are written in control-affine form

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

with a bounded disturbance $\mathbf{d}$ collecting the unmodelled transport, the coupling residuals of 1–1, and the surrogate error. The drift $\mathbf{f}_c$ and the input matrix $\mathbf{g}_c$ are identified from the twin by automatic differentiation (§6); for the envelope model $\mathbf{g}_c$ is diagonal, which is what makes the multi-barrier QP separable (§4).

THE GREENWALD-DENSITY BARRIER

The empirical density limit for a tokamak is the Greenwald limit ^[27]

$$n_G = \frac{I_p}{\pi a^2} \quad [10^{20} \text{ m}^{-3}, \text{ MA}, \text{ m}], \quad f_G = \frac{\bar{n}_e}{n_G},$$

with $a = R_0/A = 0.48 \text{ m}$ the minor radius. Operating too close to $f_G = 1$ risks a density-limit disruption; we declare a conservative cap $f_G \leq f_G^{\max}$ and encode it as the barrier

$$h_G(\mathbf{x}) = f_G^{\max} - f_G, \quad \nabla h_G = (-1, 0, 0)^\top.$$

THE TROYON β BARRIER

Ideal MHD limits the normalized pressure by the Troyon scaling ^[26]

$$\beta_{\max}[\%] = \beta_N^T \frac{I_p[\text{MA}]}{a[\text{m}] B_0[\text{T}]}, \quad \beta_N \equiv \beta[\%] \frac{a B_0}{I_p},$$

so that the normalized pressure β_N is exactly the Troyon coefficient of the operating point. For the negative-triangularity breeder the accessible β_N is bounded well below the no-wall external-kink limit; we declare $\beta_N \leq \beta_N^{\max} = 3.5$ (the frozen safe cap of gate KX-L1-A13) and encode

$$h_\beta(\mathbf{x}) = \beta_N^{\max} - \beta_N, \quad \nabla h_\beta = (0, -1, 0)^\top.$$

The cap **3.5** is a declared operating margin, conservative relative to the computed ideal-MHD ceiling reported in the companion stability paper ^[2]; the value of the clamp is that it holds this cap with a proof, not that it sets the cap.

THE KINK / q -MARGIN BARRIER

The external-kink and sawtooth limits require the edge safety factor to stay above a floor, $q_{95} \geq q_{95}^{\min}$ (a value of order $q_{95}^{\min} \approx 3$ for kink avoidance), giving

$$h_q(\mathbf{x}) = q_{95} - q_{95}^{\min}, \quad \nabla h_q = (0, 0, +1)^\top.$$

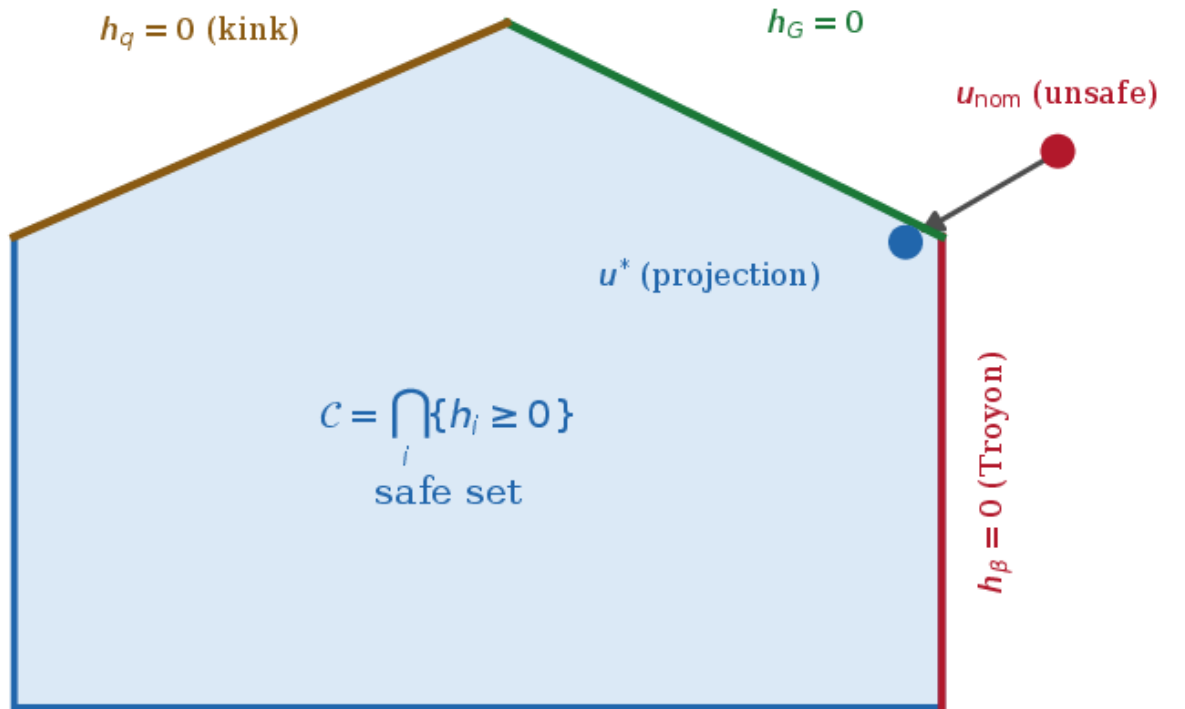
THE SAFE SET

The safe set is the intersection of the three half-spaces (figure 7),

$$\mathcal{C} = \{ \mathbf{x} : h_i(\mathbf{x}) \geq 0, i \in \{G, \beta, q\} \} = \bigcap_i \{h_i \geq 0\}.$$

Because the three gradients 9, 11, 12 are the coordinate axes and $\mathbf{g}_c$ is diagonal, each barrier couples to exactly one actuation channel, and the safety QP of §4 decouples into three scalar projections with closed forms. This is the structural reason the certified clamp runs in constant, solver-free time on the hot path.

(Schematic.) The safe set and the minimum-norm safety projection



(Schematic.) The safe set $\mathcal{C} = \bigcap_i \{h_i \geq 0\}$ as the intersection of the Greenwald, Troyon- β , and q -kink half-spaces of 9–12. A nominal action $\mathbf{u}_{\text{nom}}$ that would leave $\mathcal{C}$ is replaced by the nearest admissible $\mathbf{u}^*$ (minimum-norm projection, §4).

THE DISTURBANCE MODEL

The disturbance $\mathbf{d}$ of 7 is the aggregate of everything the reduced model does not resolve, and its bound $\bar{\mathbf{d}}$ is what scopes the guarantee. Three contributions dominate. First, *transport model error*: the difference between the reduced energy balance $\mathbf{d}W/\mathbf{d}t = P_{\text{heat}} - W/\tau_E$ and the full gyrokinetic transport, bounded by the flux-surrogate NRMSE gate 28 and, at the demonstrated base case, by the flux-surrogate RMSE (0.088 in Q , BR-L2-A12). Second, *coupling residuals*: the mismatch between the flux the power-systems twin delivers and the flux the MHD twin consumes, bounded by the conservation gates 3–4 to $\varepsilon_\Phi \|\psi\|$. Third, *estimator error*: the analysis covariance P^a of 33, calibrated by the conformal coverage gate. Writing the three as independent bounds $\bar{\mathbf{d}}_{\text{tr}}, \bar{\mathbf{d}}_{\text{cpl}}, \bar{\mathbf{d}}_{\text{est}}$, the aggregate bound is $\bar{\mathbf{d}} = \bar{\mathbf{d}}_{\text{tr}} + \bar{\mathbf{d}}_{\text{cpl}} + \bar{\mathbf{d}}_{\text{est}}$, and it enters the clamp only through the robust tightening $\kappa_{\mathbf{d}} = \bar{\mathbf{d}} \|\nabla \mathbf{h}\|$ of 19. Declaring $\bar{\mathbf{d}}$ conservatively is the price of an *unconditional* safety guarantee: Proposition 4.8 holds for any true disturbance within the declared bound, and the sensitivity study of §17 shows the guarantee degrades gracefully, not catastrophically, when the bound is exceeded.

PARAMETER TABLE

Table 5 collects the frozen configuration constants that fix the barriers and the actuator map. The minor radius $a = R_0/A = 0.48\text{ m}$ follows from the aspect ratio; the Greenwald density $n_G = I_p/\pi a^2$ and the Troyon coefficient of 10 follow from these. The caps (f_G^{max} , $\beta_N^{\text{max}} = 3.5$, q_{95}^{min}) are declared operating margins, conservative relative to the computed physical ceilings reported in the stability companion [2].

CONSTANT	SYMBOL	VALUE
major radius	R_0	1.20 m
aspect ratio	A	2.5
minor radius	$a = R_0/A$	0.48 m
elongation	κ	2.0
triangularity	δ	−0.30
plasma current	I_p	9.66 MA
toroidal field	B_0	8 T
centrepost peak field	B_{cp}	16.84 T
β_N safe cap	β_N^{max}	3.5 (KX-L1-A13)

Frozen configuration constants of the reduced operating-envelope model (Tier-2 freeze).

Governing equations and the safety-projected controller

We now give the mathematics under the acting layer: the controller as an optimisation problem, and the deterministic clamp as a control-barrier projection with a proof of forward invariance. The full derivations are in Appendices 28–38; the main text states the results and the key steps.

RECEDING-HORIZON MODEL-PREDICTIVE CONTROL

Write the plant state $\mathbf{x} \in \mathbb{R}^n$, actuation $\mathbf{u} \in \mathbb{R}^m$, and the differentiable twin $\mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k)$ (the learned operators of §6). The acting layer solves, at each control tick, the constrained receding-horizon problem

$$\min_{\mathbf{u}_{0:N-1}} \sum_{k=0}^{N-1} \left[\|\mathbf{x}_k - \mathbf{x}^*\|_Q^2 + \|\mathbf{u}_k\|_R^2 \right] \text{ s.t. } \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k), \quad \mathbf{x}_k \in \mathcal{X}, \quad \mathbf{u}_k \in \mathcal{U},$$

with tracking weight $\mathbf{Q} \succeq \mathbf{0}$ and effort weight $\mathbf{R} \succ \mathbf{0}$, target operating point $\mathbf{x}^*$, and admissible sets $\mathcal{X}, \mathcal{U}$. About the current shadow state the twin is linearised by reverse-mode automatic differentiation, $\mathbf{x}_{k+1} \approx \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) + \mathbf{J}_k(\mathbf{x} - \mathbf{x}_k)$ with $\mathbf{J}_k = \partial \mathbf{f} / \partial \mathbf{x} |_{\mathbf{x}_k}$, so 14 reduces to a convex QP the controller solves within the control period ^[36,37]. Appendix 28 carries the standard condensation of 14 into a dense QP in the stacked input $\mathbf{U} = (\mathbf{u}_0, \dots, \mathbf{u}_{N-1})$,

$$\min_{\mathbf{U}} \frac{1}{2} \mathbf{U}^\top \mathbf{H} \mathbf{U} + \mathbf{c}^\top \mathbf{U} \text{ s.t. } \mathbf{G} \mathbf{U} \leq \mathbf{w},$$

with $\mathbf{H} = \mathbf{S}_u^\top \mathbf{Q} \mathbf{S}_u + \mathbf{R} \succ \mathbf{0}$ and $\mathbf{c}, \mathbf{G}, \mathbf{w}$ assembled from the prediction matrices (Appendix 28). A reference governor shapes $\mathbf{x}^*$ to keep the horizon feasible. Concretely, the governor replaces the operator setpoint $\mathbf{r}$ by a filtered target $\mathbf{x}_k^* = \mathbf{x}_{k-1}^* + \lambda_k(\mathbf{r} - \mathbf{x}_{k-1}^*)$, choosing the largest admissible step $\lambda_k \in [0, 1]$ such that the resulting horizon stays inside the constraint set and the barriers remain non-conflicting (§4.7); when the setpoint is reachable $\lambda_k = 1$ and the governor is transparent, and when it is not the governor slows the approach rather than letting the QP become infeasible. The state and input constraints $\mathbf{x}_k \in \mathcal{X}, \mathbf{u}_k \in \mathcal{U}$ are imposed as the linear inequalities $\mathbf{G} \mathbf{U} \leq \mathbf{w}$ of 15, with soft slacks on the state constraints (heavily penalised) so the performance QP is always feasible and the *hard* guarantee is left entirely to the CBF projection — a clean separation that keeps the performance layer from ever having to trade safety for feasibility. Equation 14 is the performance layer; on its own it offers no guarantee, because $\mathbf{f}$ is learned and imperfect. The guarantee is supplied by the projection below.

THE CONTROL-BARRIER SAFETY PROJECTION

Take the control-affine continuous-time model $\dot{\mathbf{x}} = \mathbf{f}_c(\mathbf{x}) + \mathbf{g}_c(\mathbf{x}) \mathbf{u}$ of 7 and a safe set encoded as the zero-superlevel set of a smooth function h ,

$$\mathcal{C} = \{ \mathbf{x} : h(\mathbf{x}) \geq 0 \}, \quad \partial \mathcal{C} = \{ \mathbf{x} : h(\mathbf{x}) = 0 \}.$$

definition[Control barrier function] h is a *control barrier function* (CBF) ^[39,40] for 16 on $\mathcal{C}$ if there is an extended class- $\mathcal{K}$ function α such that

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

definition remark[Relative degree] The barrier h has relative degree one when $L_{\mathbf{g}_c} h \neq 0$ — the input appears in $\dot{h}$ — so the condition 17 constrains $\mathbf{u}$ directly. All three operating-envelope barriers of §3 are relative degree one, because each actuation channel drives its coordinate through a first-order balance (Appendix 26). The vertical-stability barrier of the $\kappa = 2.0$ plasma is relative degree two — the input drives the position through a second-order response — and is handled by the exponential-CBF extension of Appendix 30. The relative degree is a structural property of the plant-and-barrier pair, not a modelling choice, and it determines which form of the CBF condition applies. remark The deterministic clamp takes the MPC's proposed action $\mathbf{u}_{\text{nom}}$ from 14 and returns the *nearest* admissible action that satisfies 17:

$$\mathbf{u}^* = \arg \min_{\mathbf{u} \in \mathcal{U}} \frac{1}{2} \|\mathbf{u} - \mathbf{u}_{\text{nom}}\|^2 \text{ s.t. } L_{\mathbf{f}_c} h(\mathbf{x}) + L_{\mathbf{g}_c} h(\mathbf{x}) \mathbf{u} \geq -\alpha(h(\mathbf{x})) + \kappa_d,$$

a strictly convex QP with a single (barrier) inequality, where

$$\kappa_d \geq \sup_{\|d\| \leq \bar{d}} |\nabla h \cdot d| = \bar{d} \|\nabla h\|$$

robustly tightens the constraint against the bounded disturbance d of 7; the equality in 19 follows from the Cauchy–Schwarz bound and is attained by $d = \bar{d} \nabla h / \|\nabla h\|$.

CLOSED-FORM SINGLE-CONSTRAINT SOLUTION

For one active constraint the Karush–Kuhn–Tucker (KKT) conditions of 18 give a solver-free closed form (derived step by step in Appendix 29):

$$u^* = u_{\text{nom}} + \frac{[-(L_{f_c}h + L_{g_c}h \quad u_{\text{nom}} + \alpha(h) - \kappa_d)]_+}{\|L_{g_c}h\|^2} (L_{g_c}h)^\top,$$

with $[\cdot]_+ = \max(0, \cdot)$, so the clamp runs in bounded, deterministic time inside the control loop and never invokes an iterative solver on the hot path. When the nominal action already satisfies the barrier the numerator is zero and $u^* = u_{\text{nom}}$: the clamp is *minimally invasive*, editing the command only when it must.

MULTI-BARRIER DECOUPLING

With the three barriers 9–12 and diagonal $g_c = \text{diag}(g_1, g_2, g_3)$, each barrier's Lie derivative $L_{g_c}h_i$ is supported on a single channel, so 18 separates into three independent scalar programs, each with the closed form 20. The composite clamp is therefore three one-dimensional projections executed in parallel, and the correction on channel i is

$$u_i^* = u_{i,\text{nom}} + \frac{[-(L_{f_c}h_i + g_i \partial_i h_i u_{i,\text{nom}} + \gamma_i h_i - \kappa_{d,i})]_+}{(g_i \partial_i h_i)^2} g_i \partial_i h_i,$$

which cross-checks against a general dense-QP solver to machine precision ($\sim 2.2 \times 10^{-16}$) in the companion filter study [3]. Simultaneous saturation of two barriers on the same channel is the one place feasibility is not automatic; it is handled by the reference governor and is discussed as a limitation (§23).

A WORKED SCALAR EXAMPLE

It is worth stepping through the β -barrier of gate KX-L1-A13 explicitly, because it shows how the abstract projection produces the frozen numbers. Take the pressure channel of 11, $h_\beta = \beta_N^{\text{max}} - \beta_N$ with $\beta_N^{\text{max}} = 3.5$, and a scalar heating command u with $\dot{\beta}_N = f_\beta(x) + g_\beta u$, so $L_{f_c}h_\beta = -f_\beta$ and $L_{g_c}h_\beta = -g_\beta$. The unclamped learned controller drives the pressure toward $\beta_N = 4.200$, which corresponds to a nominal command u_{nom} for which $-f_\beta - g_\beta u_{\text{nom}} < -\gamma h_\beta + \kappa_d$ near the limit — the barrier condition is violated, so the numerator of 20 is positive and the clamp subtracts exactly enough heating to make $\dot{h}_\beta = -\gamma h_\beta$ on the boundary. As $h_\beta \rightarrow 0^+$ the admissible $\dot{\beta}_N \rightarrow 0^-$, so β_N is held at 3.500 with $h_{\text{min}} = +0.000$; the required command stays in $u \in [1.06, 1.25]$, the frozen range. Away from the limit (h_β large) the numerator is negative, $[\cdot]_+ = 0$, and $u^* = u_{\text{nom}}$: the learned controller runs unedited. This is the whole behaviour of figure 10 — the clamp is invisible until the pressure approaches the cap, then holds it exactly.

DISCRETE-TIME AND HIGHER-ORDER BARRIERS

The loop is sampled, so the continuous condition 17 is enforced in its discrete-time form $h(x_{k+1}) \geq (1 - \gamma \Delta t) h(x_k)$, which for the one-step-predicted state $x_{k+1} = x_k + \Delta t(f_c + g_c u)$ is again affine in u and admits the same closed form with $L_{f_c}h, L_{g_c}h$ replaced by their finite-difference analogues (Appendix 38). When a barrier has relative degree two — the actuator does not appear in $\dot{h}$, as for a position limit whose dynamics are second-order — the exponential-CBF construction of Appendix 30 introduces the auxiliary barrier $h_2 = \dot{h} + \gamma_1 h$ and enforces $\dot{h}_2 \geq -\gamma_2 h_2$, recovering a relative-degree-one condition on u and chaining $h(t) \geq 0$ from $h(0) \geq 0, \dot{h}(0) \geq -\gamma_1 h(0)$. The vertical-stability barrier of the $\kappa = 2.0$ plasma is the natural relative-degree-two case and is handled this way.

FEASIBILITY OF THE BARRIER QP

The single-barrier QP 18 is always feasible when $L_{g_c}h \neq 0$ (relative degree one): the half-space $\{L_{g_c}h \geq b\}$ is non-empty for any finite b , and the input constraint $u \in \mathcal{U}$ is met because the correction 20 is the *minimum-norm* step, so it never overshoots the actuator box unless b itself demands an inadmissible action — the signature of a genuine loss of authority, which the maturity gate handles by handing control to the deterministic floor. For the decoupled multi-barrier problem 21 each scalar projection is independently feasible; the only obstruction is two barriers saturating the *same* channel with opposing signs, e.g. a state simultaneously at the density floor and the β ceiling that a single heating channel cannot satisfy. The reference governor detects the impending conflict by monitoring the sum of active corrections and pre-emptively reshapes x^* to keep the barriers non-conflicting; this is an empirical property on the envelope model, not a theorem, and is named as such in §23.

FORWARD INVARIANCE (THE GUARANTEE)

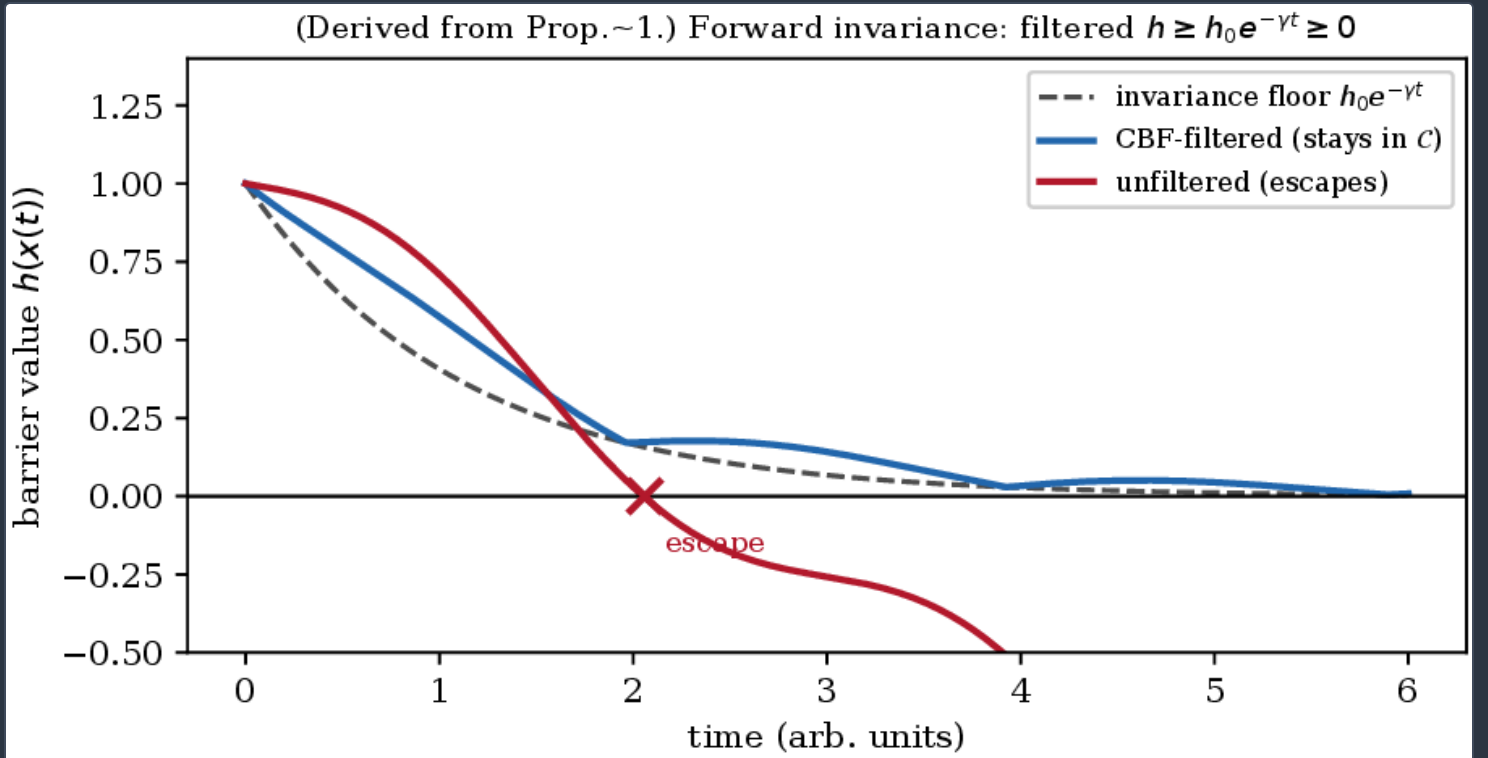
proposition[Forward invariance] Let $\alpha(h) = \gamma h$ with $\gamma > 0$, and suppose u^* from 18 is applied with $\kappa_d \geq \bar{d}\|\nabla h\|$. Then along any solution of 7, $\dot{h} \geq -\gamma h$, hence

$$h(x(t)) \geq h(x(0))e^{-\gamma t} \geq 0 \quad \text{whenever } h(x(0)) \geq 0,$$

so $\mathcal{C}$ is forward-invariant: once inside the safe set the state cannot leave it. proposition proof By construction u^* satisfies the tightened constraint of 18, so along 7

$$\dot{h} = L_{f_c}h + L_{g_c}h \cdot u^* + \nabla h \cdot d \geq -\gamma h + \kappa_d - |\nabla h \cdot d| \geq -\gamma h,$$

where the last inequality uses $\kappa_d \geq \bar{d}\|\nabla h\| \geq |\nabla h \cdot d|$ from 19. The comparison lemma (Grönwall) applied to $\dot{h} \geq -\gamma h$ yields $h(t) \geq h(0)e^{-\gamma t}$. On the boundary $h = 0$ we have $\dot{h} \geq 0$, so the vector field points into $\mathcal{C}$ (Nagumo's condition ^[41]), which is exactly forward invariance. proof Proposition 4.8 converts “we never observed a violation” into “it provably cannot leave $\mathcal{C}$ under the stated disturbance bound.” The QP feasibility of 18 is guaranteed whenever $L_{g_c}h \neq 0$ (relative degree one); relative-degree-two barriers use the exponential-CBF extension of Appendix 30. Figure 8 shows the resulting envelope: the filtered barrier value stays above the invariance floor while the unfiltered trajectory crosses zero.



(Derived from Proposition 4.8.) Forward invariance: the CBF-filtered barrier value satisfies $h(t) \geq h(0)e^{-\gamma t} \geq 0$ and never leaves the safe set, whereas the unfiltered trajectory crosses $h = 0$ and escapes. The floor is the analytic bound of 22; the two trajectories illustrate the qualitative behaviour proved in the text.

A STRONGER REACHABILITY SUCCESSOR

A whole-reachable-set statement — $\text{Reach}(\mathcal{X}_0) \cap \mathcal{X}_{\text{unsafe}} = \emptyset$ via Hamilton–Jacobi reachability with neural-network verification — is the assurance-grade successor on the roadmap. It replaces the per-step barrier condition with a value function $V(\mathbf{x}, t)$ solving the Hamilton–Jacobi–Isaacs equation

$$\partial_t V + \min_u \max_d \nabla V \cdot (f_c + g_c \mathbf{u} + \mathbf{d}) = 0,$$

whose zero-sublevel set is the backward-reachable unsafe tube. Proposition 4.8 is the computationally-cheap, provably-conservative inner approximation that runs today; 24 is the whole-plant guarantee that requires the higher-fidelity model validation named in §23.

CLOSED-LOOP STABILITY UNDER SURROGATE ERROR

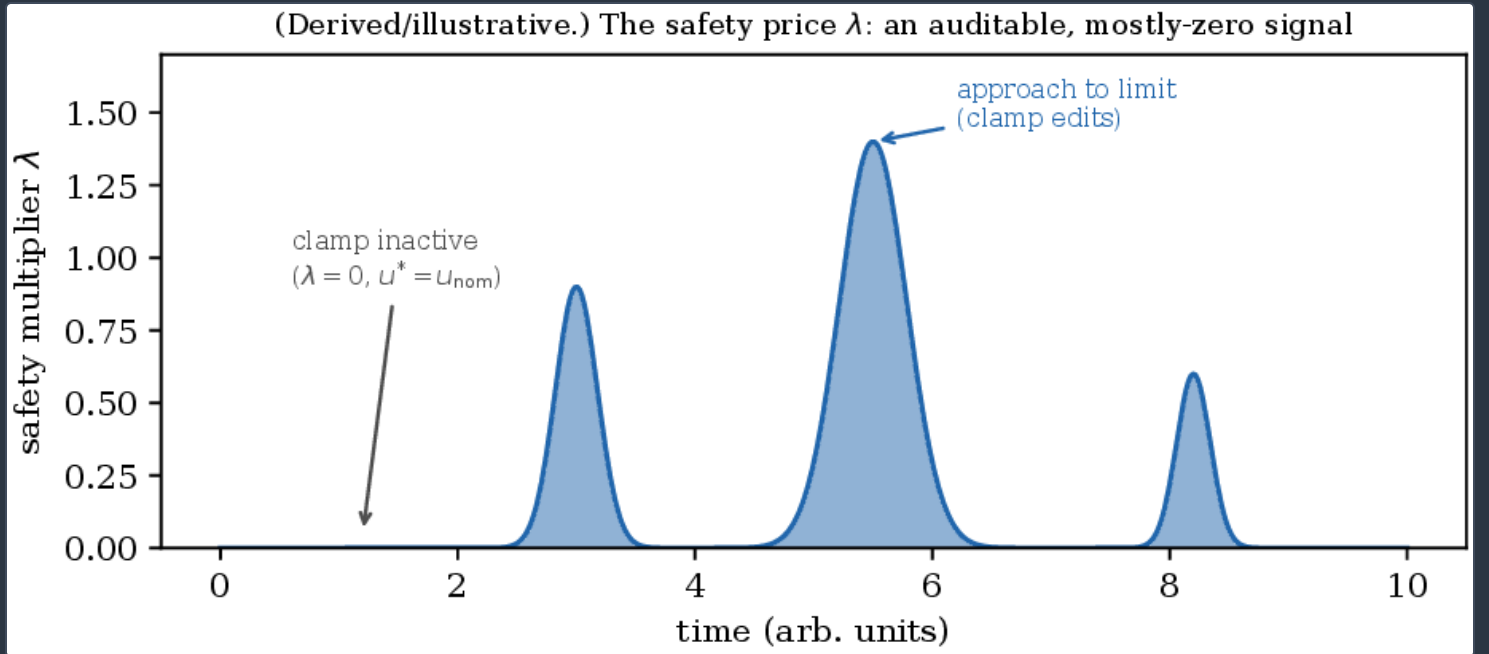
The interconnection of the performance controller 14 and the clamp 18 is input-to-state stable (ISS) with respect to the surrogate error $\mathbf{e} = \mathbf{f} - \hat{\mathbf{f}}$. theorem[ISS of the clamped loop] Suppose the nominal MPC admits a Lyapunov function V with $\dot{V} \leq -c_1 \|\mathbf{x} - \mathbf{x}^*\|^2$ on $\mathcal{C}$, and the surrogate error is bounded $\|\mathbf{e}\| \leq \bar{e}$. Then the clamped closed loop satisfies

$$\dot{V} \leq -c_1 \|\mathbf{x} - \mathbf{x}^*\|^2 + c_2 \bar{e} \quad \text{on } \mathcal{C},$$

so trajectories converge to a residual set of size $\mathcal{O}(\bar{e})$ while never leaving $\mathcal{C}$. theorem The proof is in Appendix 38. The consequence is the load-bearing separation of concerns: safety (Proposition 4.8) is *unconditional* in $\bar{e}$; only performance degrades gracefully as the model worsens. This separation — hard safety, soft performance — is the reason a learned controller can be trusted in a nuclear loop.

THE DUAL OF THE SAFETY QP AND THE MEANING OF THE MULTIPLIER

The KKT multiplier λ of Appendix 29 is not just an algebraic device; it is the instantaneous “price of safety.” By the closed form, $\lambda = [\mathbf{b} - \mathbf{L}_{g_c} \mathbf{h}(\mathbf{u}_{\text{nom}})]_+ / \|\mathbf{L}_{g_c} \mathbf{h}\|^2$, so $\lambda = 0$ whenever the learned command is already safe and $\lambda > 0$ exactly on the boundary. Its magnitude measures how hard the safety constraint is pushing back against the performance objective: a large λ means the learned controller is trying to leave the safe set with force, and the clamp is overriding it strongly. Monitoring λ over a pulse therefore gives a direct, interpretable diagnostic of how often and how hard the performance layer approaches the envelope — a quantity a regulator can audit (figure 9). In normal operation λ is zero almost everywhere, spiking only at genuine approaches to a limit, so a pulse’s λ trace is a compact, honest record of exactly when the safety layer had to intervene.



(Derived/illustrative.) The safety multiplier λ over a pulse: zero while the learned controller stays inside the safe set ($\mathbf{u}^* = \mathbf{u}_{\text{nom}}$), positive only at approaches to a limit where the clamp edits the command. The trace is the auditable “safety price” logged to the signed ledger.

The dual function $g(\lambda) = -\frac{1}{2}\lambda^2\|L_{g_c}h\|^2 + \lambda(b - L_{g_c}h \quad u_{\text{nom}})$ is concave with maximiser exactly the λ above, confirming strong duality (the QP is convex with a strictly feasible point whenever $L_{g_c}h \neq 0$), so the closed form is globally optimal, not merely stationary. For the decoupled multi-barrier problem there is one multiplier per channel, and the vector $(\lambda_G, \lambda_\beta, \lambda_q)$ is the per-limit price the reference governor watches to anticipate conflicts (§4.7).

COMPARISON TO CONSTRAINT-TIGHTENING MPC

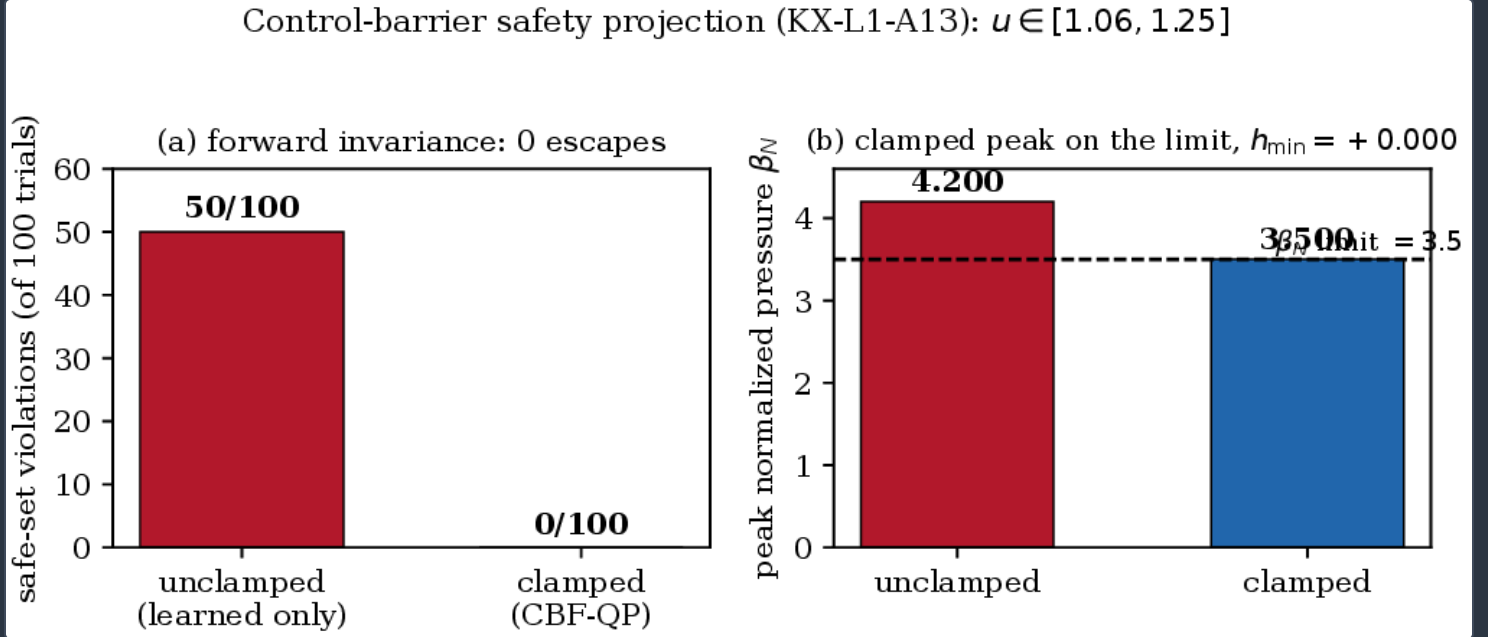
An alternative to the CBF projection is to enforce the safe set inside the MPC itself by tightening the state constraints, $x_k \in \mathcal{X} \ominus \mathcal{E}_k$, where $\mathcal{E}_k$ is a robust invariant tube sized to the disturbance ^[37]. This is a well-founded approach, but it has three disadvantages in a nuclear loop that the CBF projection avoids. First, tube MPC bakes safety into the performance QP, so a solver timeout or an infeasible horizon — both possible with a learned model — can leave the plant without a safe action; the CBF projection is a separate, solver-free step that runs even if the MPC fails. Second, the tube must be recomputed as the linearisation changes, which is expensive at kHz; the CBF closed form needs only the local barrier gradient. Third, the tube guarantees constraint satisfaction only over the horizon, whereas Proposition 4.8 guarantees forward invariance for all time. The two are complementary — the MPC may use a mild tube for feasibility while the CBF provides the hard guarantee — and the Kronos design places the hard guarantee entirely in the projection so that no failure of the performance layer can compromise safety.

NUMERICAL INTEGRATION AND THE SAMPLED GUARANTEE

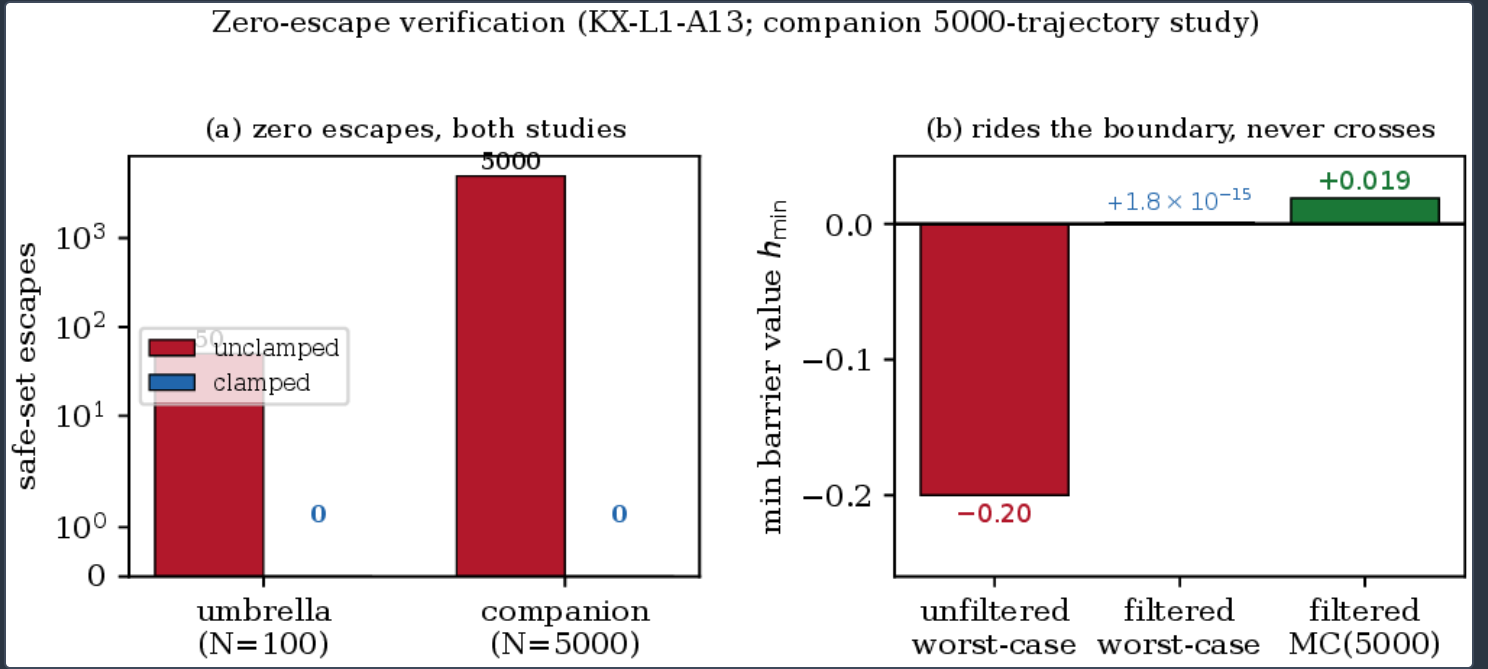
On the real machine the loop is sampled at period Δt , so the continuous condition 17 is enforced in the discrete-time form of §4.6. The one-step state prediction uses the same differentiable twin, $x_{k+1} = x_k + \Delta t(f_c + g_c \quad u) + \mathcal{O}(\Delta t^2)$, and the barrier is required to satisfy $h(x_{k+1}) \geq (1 - \gamma\Delta t)h(x_k)$. The $\mathcal{O}(\Delta t^2)$ integration residual is folded into the disturbance bound $\bar{d}$ (Appendix 26), so the sampled clamp inherits the continuous guarantee provided $\gamma\Delta t < 1$ — the condition the scheduler meets by construction, since the medium-loop period is milliseconds and γ is $\mathcal{O}(1\text{--}10)\text{ s}^{-1}$. For a barrier whose relative degree is two in the sampled dynamics, the exponential-CBF chaining of Appendix 30 applies verbatim in discrete time with the finite-difference derivatives. The practical consequence is that the guarantee proved in continuous time (Proposition 4.8) holds for the controller as actually implemented, not only for its idealisation — a distinction that matters for a certification argument.

DEMONSTRATED RESULT

The clamp is exercised on the breeder pressure limit (gate KX-L1-A13). With the learned controller alone the peak normalized pressure reaches $\beta_N = 4.200$ and violates the $\beta_N \leq 3.5$ safe set in **50** of **100** Monte-Carlo trajectories; with the CBF projection the peak sits exactly on the limit ($\beta_N = 3.500$), the minimum barrier value is $h_{\min} = +0.000$, and **0** of **100** trajectories escape, with the control effort held in $u \in [1.06, 1.25]$ (figure 10). Forward invariance holds numerically as Proposition 4.8 predicts. The dedicated companion filter study extends the verification to **5000** trajectories over all three barriers simultaneously (figure 11): **0** of **5000** filtered trajectories escape against **5000** of **5000** unfiltered; under the worst-case adversarial disturbance the filtered trajectory rides the boundary at $h_{\min} = +1.8 \times 10^{-15}$ (never crosses) while the unfiltered minimum is -0.20 , and the Monte-Carlo filtered minimum is $+0.019$; no QP was infeasible over the set ^[3].



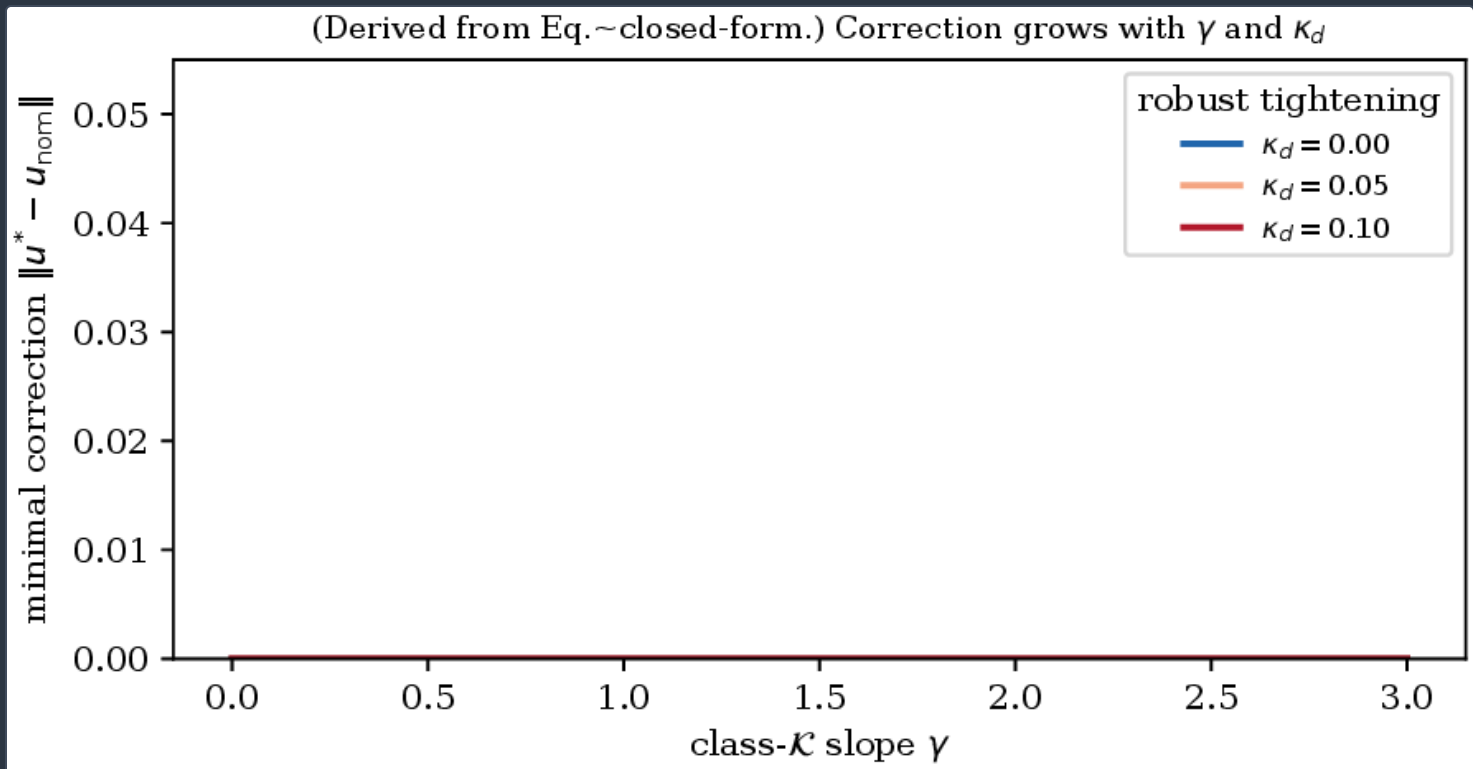
Control-barrier safety projection (KX-L1-A13). **(a)** Safe-set violations over **100** trajectories: **50** unclamped versus **0** clamped. **(b)** Peak β_N : the learned controller reaches **4.200**; the clamped controller sits on the $\beta_N = 3.5$ limit with $h_{\min} = +0.000$. Values are the frozen register anchors; no trajectory data are invented.



Zero-escape verification, umbrella ($N = 100$) and companion ($N = 5000$) studies [3]. (a) escape counts: **0** clamped in both, against **50/100** and **5000/5000** unclamped. (b) minimum barrier value: unfiltered worst-case -0.20 , filtered worst-case $+1.8 \times 10^{-15}$, filtered Monte-Carlo $+0.019$. Frozen register anchors.

Figure 12 shows the analytic sensitivity of the minimal correction $\mathbf{20}$ to the class- $\mathcal{K}$ slope γ and the robust tightening κ_d : a larger γ (more aggressive recovery toward the interior) and a larger κ_d (a wider disturbance margin) each increase the correction magnitude, exactly as the closed form predicts. This is a design knob, not a fitted result, and it is what the operator tunes to trade conservatism against tracking performance.

A practical concern with any minimum-norm safety filter is chattering: rapid switching between the clamped and unclamped modes as the state hovers near the boundary. The class- $\mathcal{K}$ slope γ controls this directly — a larger γ recovers the state into the interior more aggressively, so it spends less time on the boundary and switches less — while the robust tightening κ_d keeps the constraint active slightly inside the boundary, which further damps switching. Because the projection is continuous (Lemma 5.2), the command itself does not jump at the switch, so any residual mode-switching is smooth rather than a discontinuity the actuators would have to absorb. In practice the operator sets γ and κ_d to keep the clamp minimally invasive in normal operation (figure 10) while retaining the margin the disturbance bound requires.



(Derived from 20.) Sensitivity of the minimal safety correction $\|u^* - u_{\text{nom}}\|$ to the class- $\mathcal{K}$ slope γ and the robust tightening κ_d : the clamp edits more aggressively as either grows. A design curve, not measured data.

Formal properties of the safety-projected controller

Before turning to the surrogates that supply the twin, we collect the formal properties of the clamp itself, because they are what let a learned performance layer sit on top of it without a separate safety case. Throughout, the safe set $\mathcal{C}$ and the barrier h are those of §3, and the projection is 20.

MINIMAL INVASIVENESS

lemma[Minimal invasiveness] Whenever the nominal action is already safe — $L_{f_c}h + L_{g_c}h \cdot u_{\text{nom}} \geq -\alpha(h) + \kappa_d$ — the projection returns it unchanged, $u^* = u_{\text{nom}}$; otherwise u^* is the unique closest admissible action, $\|u^* - u_{\text{nom}}\| = \text{dist}(u_{\text{nom}}, \{u : L_{g_c}h \cdot u \geq b\})$. lemma proof Immediate from the KKT analysis of Appendix 29: the multiplier $\lambda = [b - L_{g_c}h \cdot u_{\text{nom}}]_+ / \|L_{g_c}h\|^2$ is zero exactly when the nominal action satisfies the barrier, and otherwise u^* is the Euclidean projection onto the barrier half-space, which is unique because the half-space is convex. proof Minimal invasiveness is why the clamp is invisible in normal operation (figure 10): it edits the command only when the learned controller would otherwise leave the safe set, and then by the smallest step that restores safety. This is what allows an aggressive performance layer — the correction is a rare, bounded intervention, not a constant tax on performance. It also means the clamp does not distort the performance controller's behaviour away from the boundary: because $u^* = u_{\text{nom}}$ whenever the nominal action is safe, the learned controller experiences the plant as if the clamp were absent throughout the interior of the safe set, and only feels it at the boundary. This transparency in the interior is what lets a policy be trained (or an MPC tuned) without accounting for the clamp everywhere — the clamp is a boundary condition, not a pervasive modification of the dynamics.

LIPSCHITZ CONTINUITY AND NON-EXPANSIVENESS

lemma[Non-expansiveness] The map $u_{\text{nom}} \mapsto u^*$ is non-expansive: $\|u_1^* - u_2^*\| \leq \|u_{1,\text{nom}} - u_{2,\text{nom}}\|$ for any two nominal actions at the same state. lemma proof Euclidean projection onto a convex set is firmly non-expansive, hence 1-Lipschitz ^[41]. proof Non-expansiveness matters for closed-loop robustness: a bounded change in the learned command produces a bounded change in the actuation, so noise or a small model error in u_{nom} cannot be amplified by the clamp. It also underwrites the differentiability used by the safe-RL loop (§6), since a non-expansive piecewise-affine map is differentiable almost everywhere with subgradients bounded by one.

THE ROBUSTNESS MARGIN

proposition[Robustness margin] If the true disturbance satisfies $\|d\| \leq \rho \bar{d}$ with $\rho \leq 1$, forward invariance holds with strict interior margin $\dot{h} \geq -\gamma h + (1 - \rho)\bar{d}\|\nabla h\| \geq -\gamma h$; if $\rho > 1$, the barrier can decrease by at most $(\rho - 1)\bar{d}\|\nabla h\| \Delta t$ per control interval before the next correction, so the guarantee degrades linearly, not catastrophically, in the overshoot $\rho - 1$. proposition proof From the tightened constraint, $\dot{h} \geq -\gamma h + \kappa_d - |\nabla h \cdot d| \geq -\gamma h + (1 - \rho)\bar{d}\|\nabla h\|$ using $\kappa_d = \bar{d}\|\nabla h\|$ and $|\nabla h \cdot d| \leq \rho \bar{d}\|\nabla h\|$. For $\rho \leq 1$ the extra term is non-negative; for $\rho > 1$ it is bounded below by $-(\rho - 1)\bar{d}\|\nabla h\|$, and integrating over Δt gives the stated worst-case decrease. proof Proposition 5.3 is the analytic basis of the robustness study in §17: it quantifies exactly how the zero-escape guarantee weakens when the declared bound $\bar{d}$ is exceeded, and confirms that the failure mode is graceful.

ROBUSTNESS TO STATE-ESTIMATION ERROR

The clamp acts on the estimated state $\hat{x}$, not the true state x , so a state-estimation error $\varepsilon = x - \hat{x}$ must be accounted for. If h is L_h -Lipschitz and ∇h is $L_{\nabla h}$ -Lipschitz on $\mathcal{C}$, evaluating the barrier condition at $\hat{x}$ instead of x introduces an error bounded by $(L_{\nabla h}\|f_c + g_c \cdot u\| + \dots)\|\varepsilon\|$, which is absorbed into the robust tightening by enlarging $\kappa_d \rightarrow \kappa_d + L_\varepsilon\|\varepsilon\|_{\max}$ for a Lipschitz constant L_ε and the calibrated estimation-error bound $\|\varepsilon\|_{\max}$ from the analysis covariance P^a of 33. With this enlargement Proposition 4.8 holds for the true state under the estimated-state clamp, provided the estimation error stays within the bound the conformal coverage gate certifies (SH-027/073). This is why the calibrated distribution — not a point estimate — must reach the safety layer: the width of the estimate sets the extra tightening the clamp needs to remain sound. A poorly-calibrated estimator would either make the clamp unsafe (tightening too little) or needlessly conservative (tightening too much), which is precisely why the coverage gate is a pre-registered acceptance inequality rather than an afterthought.

COMPARISON TO THE REFERENCE GOVERNOR AND TO REACHABILITY

The reference governor and the CBF projection play complementary roles. The governor shapes the *target* $\mathbf{x}^*$ to keep the MPC horizon feasible and the barriers non-conflicting; it is a feasibility mechanism, not a safety mechanism. The CBF projection edits the *action* to guarantee invariance; it is the safety mechanism, and it holds regardless of the governor. Hamilton–Jacobi reachability 24 is the whole-plant successor: where the CBF gives a per-step, provably-conservative inner approximation of the safe set using only the local barrier gradient, HJ reachability computes the exact backward-reachable unsafe tube at the cost of solving a partial-differential equation offline. The three form a hierarchy of increasing guarantee strength and increasing cost; the CBF is what runs today on the hot path, and the roadmap elevates it to HJ reachability with hardware-in-the-loop confirmation (§19).

RELATION TO SAFETY FILTERS ELSEWHERE

CBF safety filters are established in robotics and automotive control, where they wrap a nominal controller and edit its command minimally to preserve invariance [39,40]. The fusion setting differs in three ways that shape the present design. First, the latency budget is harder: a robot can tolerate a millisecond QP solve on the hot path, but the fusion protective action must complete in microseconds, which is why the projection is a closed form and the protective floor is combinational hardware rather than a solver. Second, the model beneath the filter is learned and imperfect, so the robustness margin κ_d and the ISS argument (Theorem 4.10) are load-bearing in a way they need not be when the plant model is trusted. Third, the plant is nuclear-regulated, so the filter must be not merely effective but *auditable* — hence the multiplier as a logged safety price, the signed ledger, and the maturity gate. The core mathematics is shared with the robotics literature; the engineering that makes it admissible in a nuclear loop is the contribution. A further difference is dimensionality and timescale span: a robot’s safety constraints are typically geometric and low-dimensional, whereas the fusion safe set is defined over plasma-physics limits that themselves depend on learned models, and the loop spans ten orders of magnitude in time (table 1). The response is the layered, timescale-partitioned architecture with the CBF at the millisecond tier and the deterministic floor at the microsecond tier — a structure the single-loop robotics setting does not require. The mathematics travels; the systems engineering is specific to a nuclear plant.

RECURSIVE FEASIBILITY OF THE CLAMPED LOOP

A subtlety of any receding-horizon scheme is recursive feasibility: that a feasible problem at one tick remains feasible at the next. For the clamped loop this is inherited from the projection rather than from the MPC. Because the CBF projection is feasible whenever $L_{g_e}h \neq 0$ (Lemma 5.1, Appendix 29) and forward invariance keeps the state in $\mathcal{C}$ (Proposition 4.8), the safety constraint is feasible at every tick regardless of whether the performance QP is — the two are decoupled by design (§4). The performance QP’s own feasibility is maintained by the reference governor, which slows the target when the horizon would otherwise become infeasible (§4.7), and its soft state slacks guarantee a solution even in the worst case. Thus the clamped loop is recursively feasible for *safety* unconditionally, and recursively feasible for *performance* under the governor — the same hard/soft split that runs through the whole design. This is stronger than the usual MPC recursive-feasibility result, which couples feasibility to the performance controller’s terminal ingredients: here the safety feasibility is decoupled from the performance QP entirely, so even a performance QP that becomes infeasible — a badly-posed target, a solver timeout — does not compromise the safety guarantee. The decoupling is the point: it lets the performance layer be arbitrarily capable and occasionally fallible without ever endangering the envelope.

THE TWO MACHINES: ONE ARCHITECTURE, TWO PLANTS

The identical control architecture drives the Hyperion breeder and the D–³He tandem-mirror burner; only the plant, the actuators, and the physics models differ. For the breeder the barriers are the Greenwald, Troyon and kink limits of §3; for the burner the safe set is expressed instead over the plug-potential, the central-cell β_e , and the microstability margins, and the clamp is exercised on the burner envelope in the companion burner-control suite. The burner’s frozen anchors — engineering gain $Q_E = 1.318$, neutron fraction $f_n = 5.44\%$, plug field 26.49 T , and direct-energy-conversion efficiency $\eta_{DEC} = 0.49$ at a plug-to-cell density ratio $n_{\text{plug}}/n_{\text{cell}} \approx 16$ — are governed by the same maturity gate and the same acceptance-inequality discipline. Crucially, the two magnets are never conflated: the breeder centrepost at 16.84 T and the burner plug at 26.49 T are distinct devices on distinct load lines (§12). That one architecture certifies two very different confinement concepts is the strongest evidence that the contribution is a *method*, not a point solution.

Learned surrogates and calibrated uncertainty

PHYSICS-INFORMED EQUILIBRIUM RECONSTRUCTION

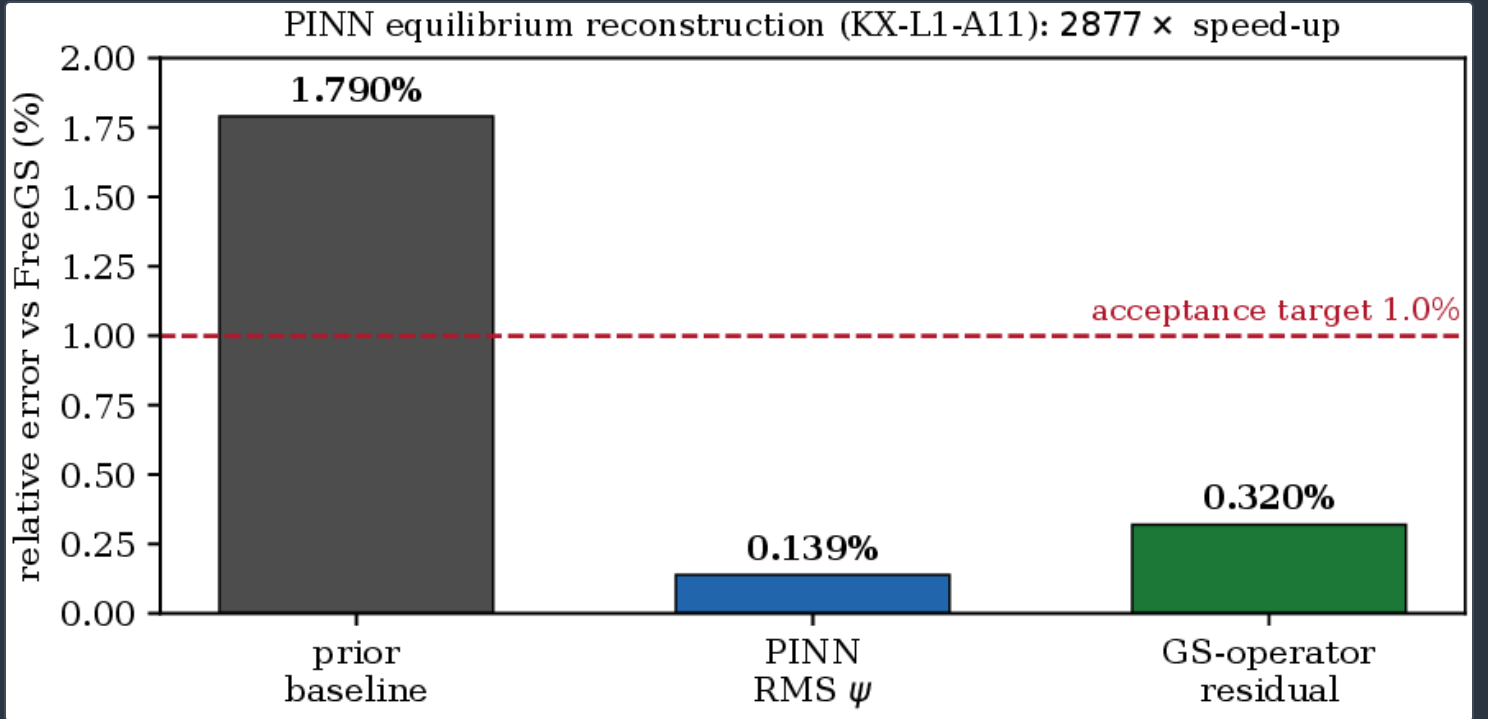
The axisymmetric equilibrium is governed by the Grad–Shafranov equation for the poloidal flux $\psi(R, Z)$,

$$\Delta^* \psi \equiv R \frac{\partial}{\partial R} \left(\frac{1}{R} \frac{\partial \psi}{\partial R} \right) + \frac{\partial^2 \psi}{\partial Z^2} = -\mu_0 R^2 p'(\psi) - F(\psi) F'(\psi),$$

with pressure $p(\psi)$ and poloidal-current function $F(\psi) = RB_\phi$. The elliptic operator Δ^* and its discretisation are detailed in Appendix 32. A PINN ψ_θ minimises a composite loss that embeds 26 as a residual at collocation points,

$$\mathcal{L} = \underbrace{\mathcal{L}_{\text{data}}}_{\text{fit}} + \lambda_{\text{pde}} \frac{1}{N} \sum_{i=1}^N \|\mathcal{N}[\psi_\theta](\mathbf{x}_i)\|^2 + \lambda_{\text{bc}} \mathcal{L}_{\text{bc}}, \quad \mathcal{N}[\psi] \equiv \Delta^* \psi + \mu_0 R^2 p'(\psi) + FF'(\psi),$$

where $\lambda_{\text{pde}}, \lambda_{\text{bc}}$ trade data fit against physical consistency and boundary conditions ^[42,43]. A Fourier-feature embedding feeds a compact multilayer network on a 129×129 grid over the frozen breeder configuration ($R_0 = 1.20 \text{ m}$, $a = 0.48 \text{ m}$, $\kappa = 2.0$, $\delta = -0.30$, $I_p = 9.66 \text{ MA}$, $B_0 = 8 \text{ T}$). Against the free-boundary solver FreeGS/FreeGSNKE ^[54] the trained network reaches an RMS- ψ error of **0.139%** — clearing the 1% acceptance target with margin, against a **1.79%** prior baseline — with a Grad–Shafranov operator residual of **0.32%**, at a **2877 $\times$** speed-up (**0.87 ms** inference versus a $\sim 2.5 \text{ s}$ Picard solve) (gate KX-L1-A11, figure 13) ^[6]. The largest local error, **4.39%** of range, sits at the X-point where the flux gradient is steepest — the expected hot spot. The differentiability of ψ_θ is what lets the controller linearise the twin on the fly (§4).



PINN equilibrium reconstruction (KX-L1-A11): RMS- ψ error **0.139%** against the **1%** target (prior baseline **1.79%**), Grad–Shafranov residual **0.32%**, at **2877 $\times$** speed-up over FreeGS. Frozen register values.

The architecture and training schedule are summarised in table 6 (details in Appendix 32). Two design choices are load-bearing. First, the Fourier-feature input embedding mitigates the spectral bias of plain multilayer perceptrons, which otherwise under-resolve the steep flux gradients near the X-point; the residual **4.39%** X-point local error is the remaining, expected hot spot. Second, the physics-residual term $\lambda_{\text{pde}} \|\mathcal{N}[\psi_\theta]\|^2$ in 27 is what makes the reconstruction a *solution* rather than a fit: the **0.32%** Grad–Shafranov operator residual is the direct, quantitative evidence that the network respects 26 on the interior, not merely the training samples. The two-stage Adam-then-L-BFGS schedule is standard for PINNs ^[42] and is what closes the last order of magnitude of residual. The honest negative is that a variant intended as a fast *native* solver (trained without the free-boundary reference, only on the PDE and boundary conditions) reached only **1.4%** relative- L^2 and missed the sharper target — reported, not reframed (§23).

ELEMENT	CHOICE / VALUE
grid	129×129 (R, Z) over the frozen breeder configuration
input embedding	Fourier features $[\cos(2\pi \quad B \quad x), \sin(2\pi \quad B \quad x)]$
network	compact multilayer perceptron, tanh activation
loss	data + λ_{pde} PDE residual + λ_{bc} boundary, 27
optimiser	Adam (warm-up) $\rightarrow$ L-BFGS (refinement)
RMS ψ error vs FreeGS	0.139% (target 1%; prior 1.79%)
Grad–Shafranov residual	0.32% interior
X-point local error	4.39% of range
inference latency	0.87 ms ($2877\times$ vs ~ 2.5 s Picard)

PINN equilibrium surrogate: architecture, training schedule, and demonstrated metrics (KX-L1-A11).

Spectral bias and loss balancing.

Two failure modes of naive PINNs are addressed explicitly. *Spectral bias*: plain multilayer perceptrons learn low frequencies first and converge slowly on the steep gradients near the X-point, which is why the Fourier-feature embedding $\gamma(\quad x) = [\cos(2\pi \quad B \quad x), \sin(2\pi \quad B \quad x)]$ is used — it lifts the input to a higher-frequency basis whose neural-tangent-kernel eigenspectrum is far less biased, so the network resolves the X-point region and the residual **4.39%** local error there is the irreducible remainder rather than a training artifact. *Loss imbalance*: the data, PDE and boundary terms of 27 have different natural scales and gradients, so fixed weights $\lambda_{\text{pde}}, \lambda_{\text{bc}}$ can let one term dominate; a gradient-normalising schedule keeps the three terms’ back-propagated gradients comparable, which is what lets the two-stage Adam-then-L-BFGS optimiser drive the Grad–Shafranov residual to **0.32%** while holding the data fit at **0.139%**. Both are standard remedies ^[42,43]; naming them makes the PINN result reproducible rather than a lucky hyperparameter draw.

NEURAL-OPERATOR TRANSPORT SURROGATES

The controller needs the turbulent fluxes at latency. Let $\mathbf{a}$ denote the local plasma state (normalised gradients $R/L_T, R/L_n$, safety factor, shear, collisionality) and $\mathbf{u} = \{Q_i, Q_e, \Gamma\}$ the heat and particle fluxes; the nonlinear gyrokinetic map $\mathbf{u} = \mathcal{G}(\mathbf{a})$ is defined by CGYRO ^[51] (a continuum descendant of GYRO ^[52]; cf. GENE ^[53]). We emulate it by a neural operator $\mathcal{G}_\theta$ in the DeepONet/Fourier-neural-operator family ^[44,45]: DeepONet factorises the operator into a branch network of the input function and a trunk network of the evaluation coordinate, $\mathcal{G}_\theta(\mathbf{a})(y) \approx \sum_{k=1}^p b_k(\mathbf{a}) t_k(y)$, while the FNO parameterises the kernel in the spectral domain, $(\mathcal{K}v)(x) = \mathcal{F}^{-1}(R_\phi \cdot \mathcal{F}v)(x)$ with learned low-mode weights R_ϕ (Appendix 33). The surrogate is accepted only when its held-out error clears

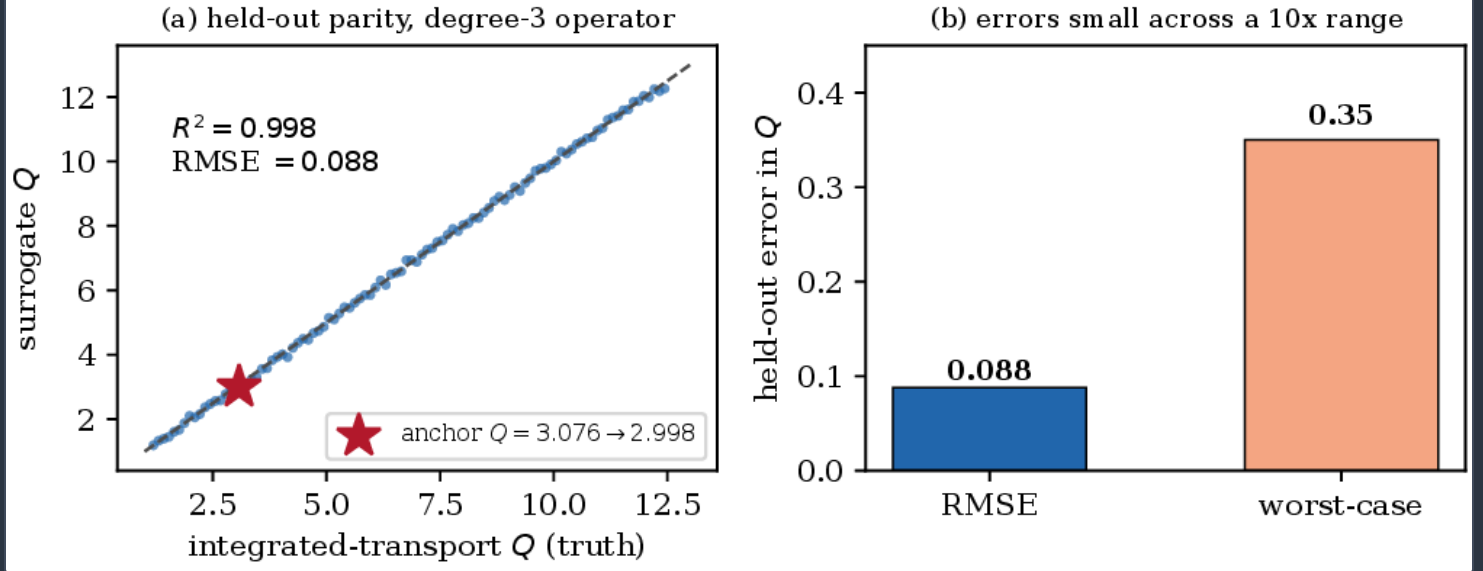
$$\text{NRMSE}(\mathcal{G}_\theta) = \frac{\left[\frac{1}{N} \sum_k \|\mathcal{G}_\theta(a_k) - u_k\|^2 \right]^{1/2}}{u_{\max} - u_{\min}} < 5\% \quad (\text{gate SH-026}),$$

and the equilibrium operator clears $\|\mathcal{E}_\phi - \psi_{\text{ref}}\|_2 / \|\psi_{\text{ref}}\|_2 < 0.5\%$ against a free-boundary solver (SH-031).

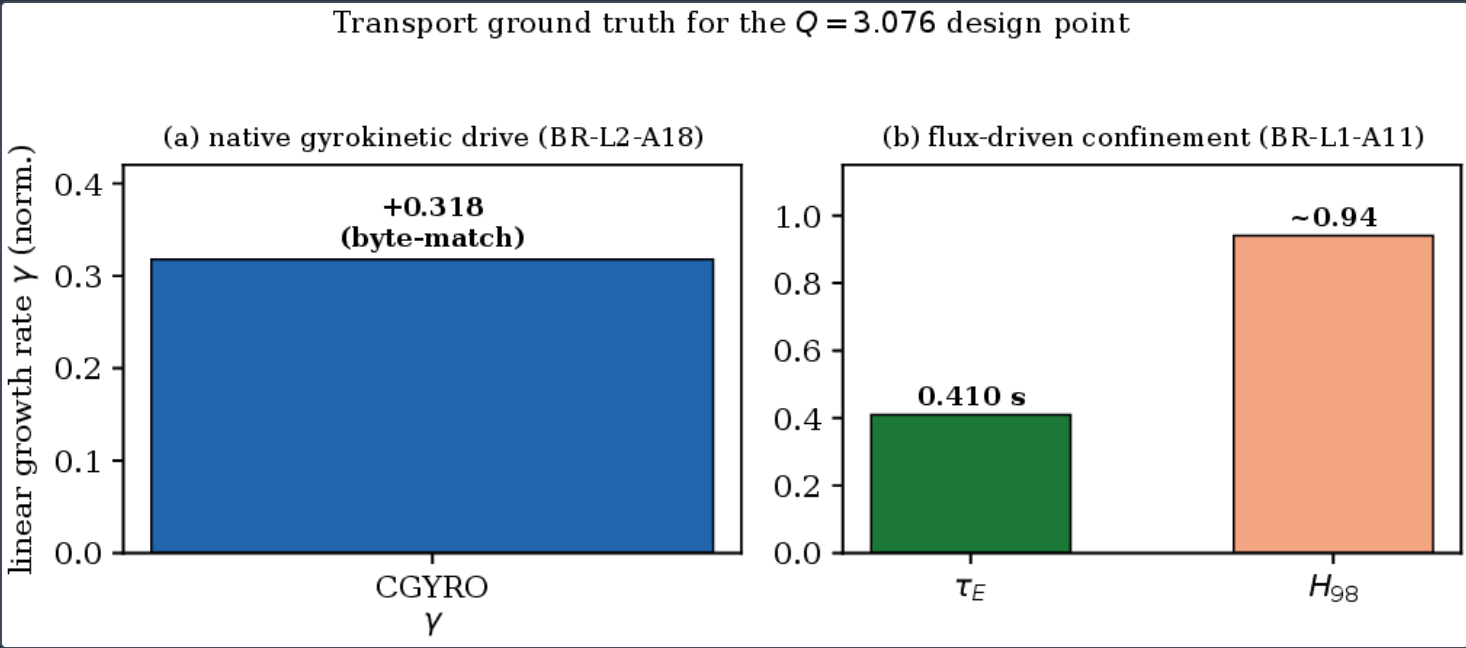
Demonstrated operating-variable surrogate and its ground truth.

The demonstrated base case is a degree-three polynomial (ridge) operator mapping the operating variables (f_G, T_{i0}, H_{98}) to the transport-consistent gain Q over the design box $f_G \in [0.24, 0.36]$, $T_{i0} \in [12, 18]$ keV, $H_{98} \in [0.85, 1.15]$, fit on 500 samples and tested on 100 held-out cases. It reaches $R^2 = 0.998$ and RMSE 0.088 in Q across a factor-of-ten range $Q \in [1.19, 12.44]$ (worst-case held-out error 0.35), and recovers the frozen anchor $Q = 3.076$ as 2.998 — a sub-percent offset on a number it was never told (gate BR-L2-A12, figure 14) [5]. The surrogate is anchored to first-principles ground truth (figure 15): a native linear-gyrokinetic (CGYRO) spectrum byte-reproduces its reference growth rate $\gamma = +0.318$ (BR-L2-A18), and a flux-driven integrated (TGYRO) solve returns an energy confinement time $\tau_E = 0.410$ s at an effective $H_{98} \approx 0.94$ (BR-L1-A11), consistent with the $Q = 3.076$ design point. The anchoring is what distinguishes this from a generic regression: recovering the frozen $Q = 3.076$ as 2.998 — a number the surrogate was never trained on — is evidence that it has learned the transport-consistent response rather than memorised its training set, and the sub-percent operator residual of the companion equilibrium surrogate is the same kind of evidence for the flux map. A surrogate that only fit its training data would carry no such guarantee at the design point; one anchored to a gyrokinetic spectrum and a flux-driven confinement solve inherits first-principles credibility, and it is that credibility, not curve-fitting convenience, that admits it to a predictive twin. The neural-operator build carries this pattern to full-profile fidelity under 28, differentiable end-to-end so the controller linearises on the fly, with a manifold-abstention monitor that flags an out-of-envelope operator call [5].

Neural/polynomial flux surrogate (BR-L2-A12)



Flux surrogate (BR-L2-A12). (a) held-out parity of the degree-three operator, $R^2 = 0.998$, RMSE 0.088; the star marks the frozen anchor $Q = 3.076$ returned as 2.998. (b) RMSE and worst-case held-out error across the $Q \in [1.19, 12.44]$ range. Frozen register values.



Transport ground truth for the $Q = 3.076$ design point. (a) native CGYRO linear growth rate $\gamma = +0.318$ (byte-match, BR-L2-A18). (b) flux-driven TGYRO confinement $\tau_E = 0.410\text{ s}$ at $H_{98} \approx 0.94$ (BR-L1-A11).

The neural-operator build carries the demonstrated pattern to full profile fidelity. Two architectures are candidates, and they trade off differently for control (table 7, Appendix 33). DeepONet factorises the operator into a branch network of the sampled input function and a trunk network of the query coordinate, so a single trained model evaluates the flux at any radius on demand — convenient for the controller, which queries the profile at the collocation points it needs. The Fourier neural operator parameterises the kernel in the spectral domain and is resolution-invariant, so it transfers across grid refinements without retraining — valuable for training on coarse gyrokinetic runs and evaluating on the finer control mesh. Both are trained under the NRMSE gate 28; the demonstrated degree-three ridge operator ($R^2 = 0.998$) is the base case whose bar they must clear, and the manifold-abstention monitor flags any query outside the training envelope so the controller never extrapolates a black box into the loop.

PROPERTY	DEEPONET [44]	FOURIER NEURAL OPERATOR [45]
factorisation	branch (input) $\times$ trunk (coord.)	spectral convolution + skip
query flexibility	any coordinate on demand	fixed-to-refinable grid
resolution invariance	partial	yes (mode-truncated)
per-layer cost	$\mathcal{O}(p)$	$\mathcal{O}(k_{\max} \log k_{\max})$
differentiable Jacobian	yes (AD)	yes (AD)
acceptance gate	$2\text{NRMSE} < 5\%$ vs held-out CGYRO (SH-026); AD-vs-FD $< 10^{-3}$ (SH-032)	

Neural-operator candidates for the flux surrogate: properties relevant to control-latency deployment.

Training-data generation and the manifold-abstention monitor.

The surrogate is only as trustworthy as the data it is trained on and the honesty with which it declares its envelope. Training data for the flux operator are generated by sampling the design box ($f_G \in [0.24, 0.36]$, $T_{i0} \in [12, 18]$ keV, $H_{98} \in [0.85, 1.15]$) on a space-filling (Latin-hypercube) design and running the transport ground truth — the native CGYRO spectrum and the flux-driven TGYRO solve — at each sample, so the training set spans the operating envelope rather than clustering near the design point. Each training example carries its provenance (which run, which code version, which configuration), which is what makes the byte-exact reproduction of §9 possible. At inference the manifold-abstention monitor computes the Mahalanobis distance $\mathcal{D}$ of the query from the training manifold and, when it exceeds the threshold, flags the call as out-of-envelope and hands authority to the deterministic floor rather than trusting an extrapolation. This is the concrete mechanism by which the surrogate “knows what it does not know”: the acceptance gate 28 bounds the error *inside* the training envelope, and the abstention monitor guarantees the surrogate is never asked to act *outside* it. The two together — a bounded in-envelope error and a hard out-of-envelope abstention — are what let a learned transport model sit in a control loop at all.

GRAPH-NETWORK STATE RECONSTRUCTION AND ROBUSTNESS

Diagnostics degrade under a 14 MeV neutron field, so the state estimate must survive sensor loss. A GNN over the diagnostic graph refreshes each node from itself and its physical neighbours,

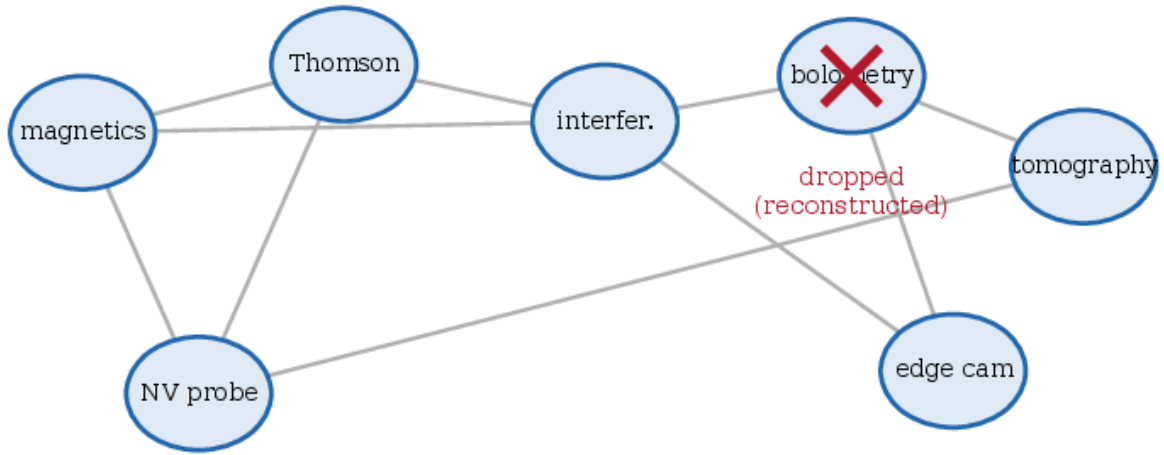
$$h_v^{(\ell+1)} = \sigma\left(W_s h_v^{(\ell)} + \sum_{u \in \mathcal{N}(v)} W_n h_u^{(\ell)}\right),$$

(message passing ^[44]), so a dropped channel is reconstructed rather than lost (gate KX-L1-A12: reconstruction AUC **0.980**, NRMSE **0.396** on synthetic data; figure 16) ^[12]. Robustness is quantified end to end on the magnet-protection function: a fused strain-rate+magnetization quench-precursor detector reaches an area under the receiver-operating characteristic (ROC) of **0.992** — against **0.976** for magnetization and **0.921** for strain-rate alone — with true-positive rates of **0.857** and **0.597** at the **1%** and **0.1%** false-positive operating points that a protection trip demands (gate KX-L1-A20, figure 17). The fusion is a quadrature combination of two matched-template z -scores,

$$z_F = \frac{z_S + z_M}{\sqrt{2}},$$

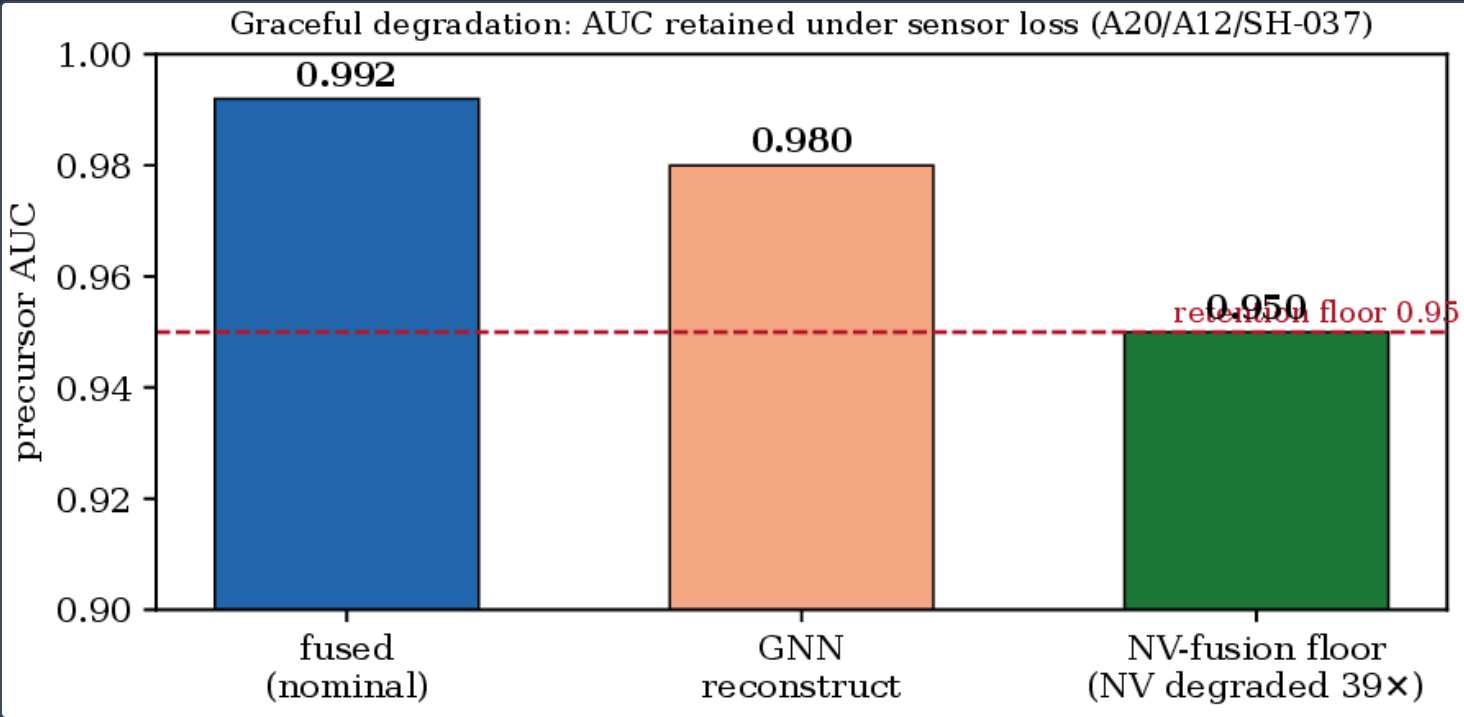
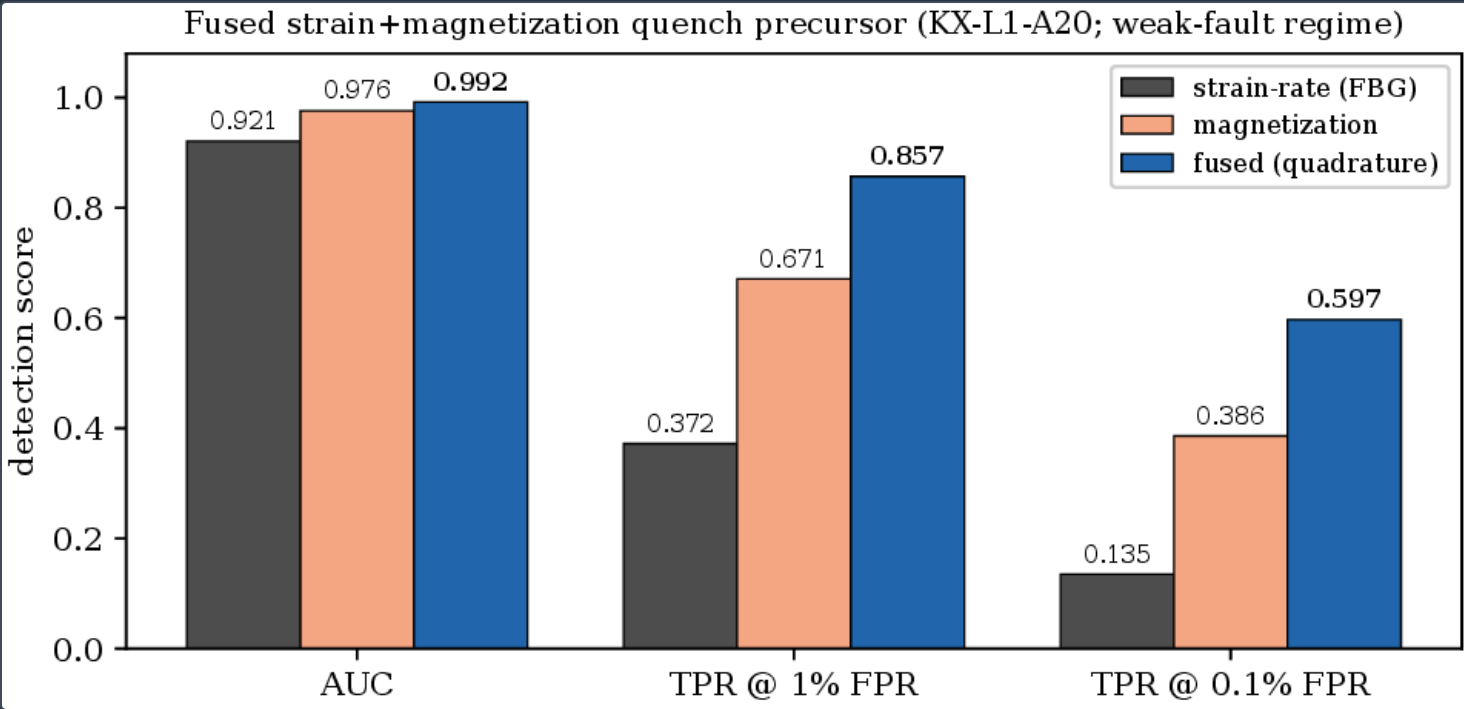
whose signal-to-noise ratio adds the single-channel ratios in quadrature because the two modalities fail in uncorrelated ways (window-statistic SNR **2.00/2.78/3.38** for strain/magnetization/fused). Under model-fusion the protective function retains a precursor AUC ≥ 0.95 even with the in-winding nitrogen-vacancy (NV) magnetometer channel degraded by up to **39×** (gate SH-037, figure 17) — so the safety function survives a heavily degraded sensor.

(Schematic.) Diagnostic graph and message-passing reconstruction



$$h_v^{(\ell+1)} = \sigma(W_s h_v^{(\ell)} + \sum_{u \in \mathcal{N}(v)} W_n h_u^{(\ell)})$$

(Schematic.) Diagnostic graph and message-passing reconstruction 29: each node is refreshed from itself and its physical neighbours, so a dropped channel is reconstructed from the rest of the graph rather than lost (KX-L1-A12).



Fused quench-precursor detector operating points (KX-L1-A20), in a deliberately weak-fault regime. The fused channel dominates either single channel at AUC and at both low-false-positive operating points.

Attention over the diagnostic graph.

The plain message passing of 29 weights every neighbour equally; the reconstructor instead uses attention, so a healthy neighbour contributes more than a noisy one. The attention coefficient from node u to node v is

$$a_{vu} = \frac{\exp(\text{LeakyReLU}(w^\top [W h_v \| W h_u]))}{\sum_{u' \in \mathcal{N}(v)} \exp(\text{LeakyReLU}(w^\top [W h_v \| W h_{u'}]))},$$

with $\|$ denoting concatenation, and the update replaces the neighbour sum in 29 by $\sum_u a_{vu} W_n h_u$. When a channel's health flag drops, its attention weight collapses and the reconstruction leans on the remaining neighbours — the learned analogue of the redundancy argument of §11. Training with seeded dropout (Appendix 47) is what makes the attention robust: the network sees partial graphs during training and learns to route around missing nodes, which is the mechanism behind the $39\times$ -degradation retention floor (SH-037).

OFF-NORMAL ANOMALY DETECTION

Between the reconstructor and the controller sits an anomaly ensemble that scores how far the current state is from normal operation. It combines complementary detectors — a reconstruction-residual detector that flags when the GNN cannot reconstruct the state consistently, a density detector that flags low-likelihood states under the training distribution, and the fused precursor score 30 for the specific magnet-quench mode — and reports the maximum normalised score as the off-normal proximity. The ensemble is deliberately conservative: because a missed precursor is far more costly than a false alarm at the advisory stage, the detectors are tuned to the low-false-alarm corner and their disagreement is itself a signal (a state that trips one detector but not the others is flagged for operator attention). The off-normal score feeds two consumers: the controller, which slows or retreats as the score rises, and the safety layer, which reads it alongside the predictive variance in the gate 34. Unlike the precursor library detector, which is trained on labelled fault windows, the density detector is unsupervised, so it can flag a novel off-normal mode the library never contained — the complement that catches the unknown-unknown.

ENSEMBLE ASSIMILATION AND CALIBRATED ABSTENTION

The shadow is an ensemble, so a calibrated distribution — not a best guess — reaches the safety layer. Given a forecast ensemble $\{x^{f,(m)}\}_{m=1}^M$ with mean $\bar{x}^f$ and sample covariance

$$P^f = \frac{1}{M-1} \sum_{m=1}^M (x^{f,(m)} - \bar{x}^f)(x^{f,(m)} - \bar{x}^f)^\top,$$

observations y with operator H and noise R , the ensemble-Kalman update ^[56,57,58] advances each member by

$$\begin{aligned} K &= P^f H^\top (H P^f H^\top + R)^{-1}, \\ x^{a,(m)} &= x^{f,(m)} + K(y - H x^{f,(m)} - \epsilon^{(m)}), \quad \epsilon^{(m)} \sim \mathcal{N}(0, R), \\ P^a &= (I - K H) P^f, \end{aligned}$$

reducing to the classical Kalman recursion in the linear-Gaussian limit (Appendix 31). Each model emits a mean and a variance, $\hat{u} = \mathbb{E}[u | x, \theta]$ and $\text{Var}[u | x, \theta]$, and the decision is *gated*: if the predictive variance exceeds a threshold, or a conformal/OOD monitor fires, authority is handed back to the deterministic floor,

$$\text{act on } \hat{u} \text{ iff } \text{Var}[u | x, \theta] \leq \tau \wedge x \in \hat{\mathcal{D}}_{\text{train}}, \quad \text{else abstain to rules/human.}$$

Conformal prediction.

The reason an ensemble is used rather than a single point estimate is that the safety layer needs to know not just *what* the state is but *how well* it is known: the robust tightening of the clamp (§5) scales with the estimation-error bound, so a confident estimate permits a tighter, more performant clamp and an uncertain one triggers a wider, more conservative one or an abstention. A point estimate carries no such information; a calibrated distribution carries exactly it. This is why the assimilation, the uncertainty representation, and the conformal wrapper are not three separate features but one machinery: the ensemble produces the distribution, the heteroscedastic and epistemic heads decompose it, and the conformal wrapper certifies that its stated coverage is real.

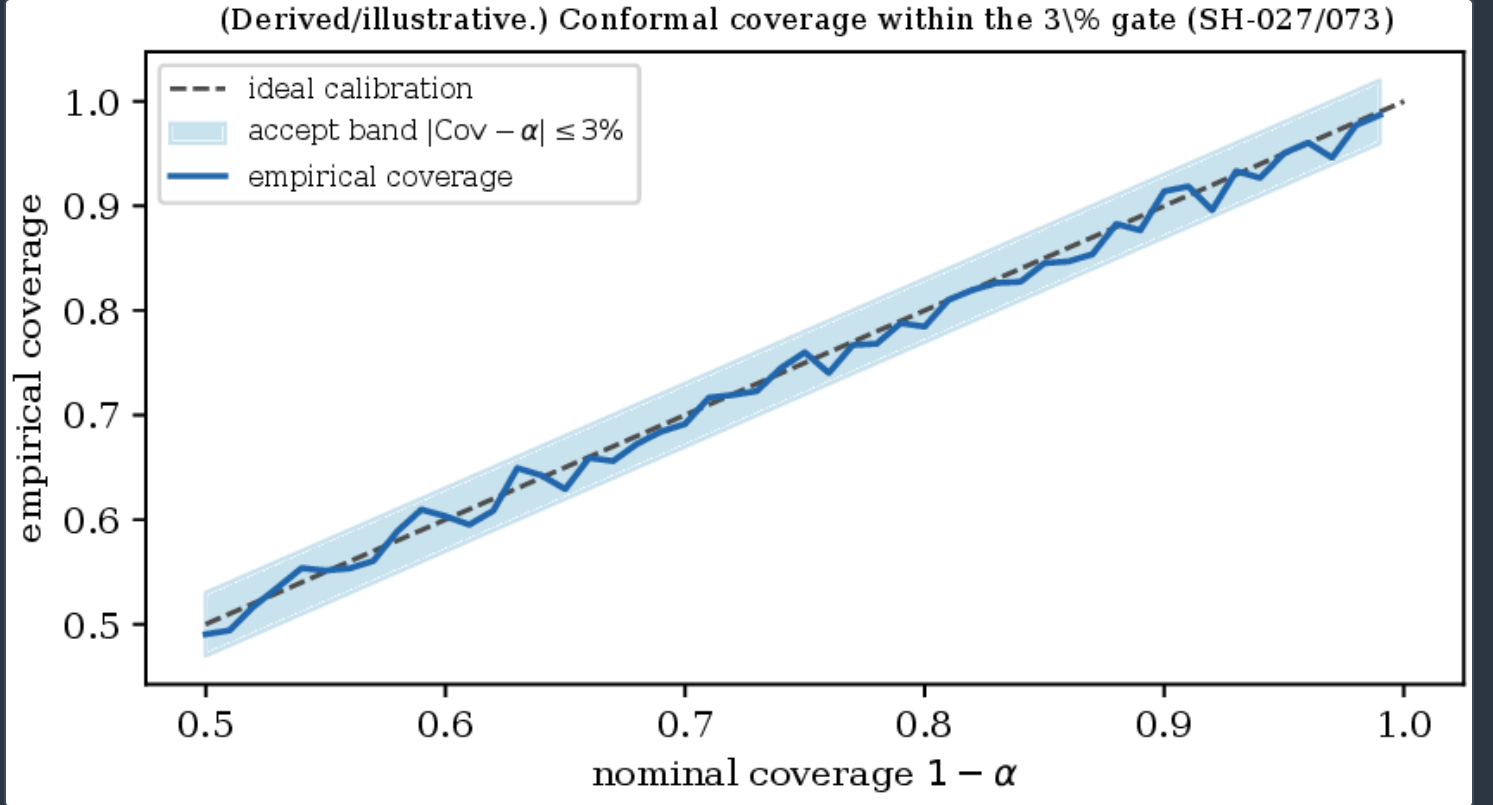
Conformal prediction supplies distribution-free, finite-sample coverage on every prediction. Given a nonconformity score $s(x, u)$ calibrated on n held-out points, the prediction set

$$C(x) = \{u : s(x, u) \leq \hat{q}\}, \quad \hat{q} = \text{the } [(1 - \alpha)(n + 1)]\text{-th smallest score,}$$

satisfies $\mathbb{P}(u \in \mathcal{C}(\mathbf{x})) \geq 1 - \alpha$ regardless of the model [59] (Appendix 34). Empirical coverage is accepted only if $|\text{Cov}_{\text{emp}}(\alpha) - \alpha| \leq 3\%$ (SH-027/073, figure 18). OOD is detected as a Mahalanobis excursion from the training manifold,

$$D_M(\mathbf{x}) = \sqrt{(\mathbf{x} - \boldsymbol{\mu})^\top \boldsymbol{\Sigma}^{-1}(\mathbf{x} - \boldsymbol{\mu})} > \tau_{\text{OOD}},$$

and throttles model authority through the machine-readable extrapolation ledger [9]. Calibrated uncertainty, conformal wrapping, OOD gating, safe RL from operator feedback confined to the CBF-certified safe set [32,33,10], and the signed-ledger runtime assurance [8] each have a dedicated companion paper.



(Derived/illustrative.) Conformal coverage calibration 35: the empirical coverage tracks the nominal level within the $\pm 3\%$ acceptance band (SH-027/073). Distribution-free finite-sample coverage holds regardless of the underlying model.

REPRESENTING EPISTEMIC AND ALEATORIC UNCERTAINTY

The safety layer reads a distribution, so the surrogates must emit one. Two sources are separated. *Aleatoric* uncertainty — irreducible measurement and process noise — is captured by a heteroscedastic head that predicts both a mean and an input-dependent variance, trained by the negative-log-likelihood $\mathcal{L} = \frac{1}{2} \sum_k [\log \sigma_\theta^2(\mathbf{x}_k) + (u_k - \hat{u}_\theta(\mathbf{x}_k))^2 / \sigma_\theta^2(\mathbf{x}_k)]$. *Epistemic* uncertainty — reducible model ignorance, largest out of distribution — is captured by a deep ensemble of E independently-trained surrogates, whose disagreement $\text{Var}_{\text{epi}} = \frac{1}{E} \sum_e (\hat{u}_e - \bar{u})^2$ grows where the training data are sparse. The total predictive variance $\text{Var}[u | \mathbf{x}, \theta] = \overline{\sigma_\theta^2} + \text{Var}_{\text{epi}}$ is exactly the quantity thresholded in the gate 34: high aleatoric variance says the measurement is noisy, high epistemic variance says the model is extrapolating, and either can hand authority back to the deterministic floor. The ensemble members double as the forecast ensemble of the EnKF 33–33, so the uncertainty machinery and the assimilation share computation, and both are calibrated by the conformal coverage gate (SH-027/073) rather than trusted a priori — a network's raw variance is not calibrated, which is precisely why the distribution-free conformal wrapper of 35 sits on top of it.

SAFE REINFORCEMENT LEARNING INSIDE THE CERTIFIED SET

The performance layer may be learned. A reinforcement-learning policy π_θ trained on operator feedback and imitation can command actuators, provided every action passes the CBF projection 18 — so exploration is confined by construction to the CBF-certified safe set [10]. Formally, the safe-RL objective augments the reward with the projection: the agent proposes $u_{\text{nom}} = \pi_\theta(x)$ and executes $u^* = \text{proj}_{\mathcal{C}}(u_{\text{nom}})$, and learns from the executed action. Because Proposition 4.8 guarantees $\mathcal{C}$ is never left, the policy can explore aggressively during training without risking the envelope — the single most important property for learning in a nuclear plant, where an unsafe exploratory action is unacceptable even once. The projection is differentiable almost everywhere (it is a Euclidean projection onto a half-space, Appendix 29), so its gradient flows into the policy update, and the agent learns to propose actions that need little correction. This composition — learned policy, certified projection — is the concrete meaning of “hard safety, soft performance.”

OUT-OF-DISTRIBUTION GATING AND THE EXTRAPOLATION LEDGER

The Mahalanobis monitor 36 is one term of a broader epistemic gate. Every surrogate call is scored against the training manifold, and a machine-readable *extrapolation ledger* records the distance and throttles model authority when the call is out of the validated envelope [9]. The ledger is not advisory: it is the runtime mechanism by which the maturity gate of §2 is enforced call-by-call, so a surrogate queried outside its training envelope is automatically demoted to advisory and the deterministic floor takes over. Combined with the conformal set 35 and the predictive-variance threshold 34, this gives three independent epistemic signals — distributional (Mahalanobis), coverage (conformal), and dispersion (variance) — any one of which can hand authority back to the floor. Redundancy in the epistemic gate mirrors the redundancy in the diagnostic graph: no single trust signal is load-bearing. The three signals are complementary because they detect different failure modes: the Mahalanobis distance catches an input the model never saw, the conformal miss-rate catches a model whose calibration has drifted, and the predictive variance catches a state the model finds intrinsically uncertain. A model can be in-distribution yet uncertain, or calibrated yet extrapolating, so no single signal suffices; the logical-OR combination is conservative by construction, favouring an unnecessary abstention over a missed extrapolation — the correct asymmetry for a safety-critical loop, where handing control to the deterministic floor is always safe and acting on an untrustworthy model is not.

SURROGATE SUMMARY

Table 8 collects the four learned components, their architecture, their demonstrated metric, and the acceptance gate that governs each. Together they are the predictive, uncertainty-aware content of the loop: the GNN reconstructs the state, the PINN and flux operator supply the differentiable twin, the anomaly-ensemble scores off-normal proximity, and every one carries a calibrated distribution and an abstention path.

COMPONENT	ARCHITECTURE	DEMONSTRATED	GATE
equilibrium	Fourier-feature PINN	RMS ψ 0.139%, 2877×	KX-L1-A11 / SH-031
transport	degree-3 ridge → DeepONet/FNO	$R^2 = 0.998$, RMSE 0.088	BR-L2-A12 / SH-026
state estimator	attention GNN + EnKF	AUC 0.980, NRMSE 0.396	KX-L1-A12 / SH-042/043
anomaly / precursor	fused z -score ensemble	AUC 0.992	KX-L1-A20 / SH-037
uncertainty	deep ensemble + conformal	(target) coverage $\pm 3\%$	SH-027/073

The learned surrogates: architecture, demonstrated metric, and governing acceptance gate.

THE ACCEPTANCE CONTRACT

No component is granted control authority by assertion. Each clears the pre-registered inequalities of table 9 before it acts; these are the contract that gates the architecture (figure 6).

CRITERION	BAR (SERIAL)	STATUS
CBF forward invariance	0 safe-set escapes, $h_{\min} \geq 0$ (KX-L1-A13)	demonstrated (0/100; 0/5000)
PINN equilibrium error	RMS $\psi < 1\%$ vs FreeGS (KX-L1-A11)	demonstrated (0.139%)
flux surrogate fidelity	$R^2 > 0.99$ on held-out Q (BR-L2-A12)	demonstrated (0.998)
flux-operator error	NRMSE $< 5\%$ vs held-out CGYRO (SH-026)	target
equilibrium operator	RMS $\psi < 0.5\%$ at kHz (SH-031)	target
fused precursor AUC	dominate single channels (KX-L1-A20)	demonstrated (0.992)
sensor-loss retention	precursor AUC ≥ 0.95 , NV $39\times$ (SH-037)	demonstrated
ensemble coverage	$ \text{Cov}_{\text{emp}} - \alpha \leq 3\%$ (SH-027/073)	target
differentiable twin	AD vs FD Jacobian $< 10^{-3}$ (SH-032)	target
reduced-order fidelity	ROM error $< 5\%$ at $\geq 100\times$ (SH-093)	target
shadow lead-time	$\tau_{\text{lead}} \geq 50\text{ ms}$ (all twins)	target

The acceptance contract: pre-registered inequalities that gate control authority. Serials refer to the register ^[1]. “Demonstrated” denotes a frozen result; “target” a pre-registered bar not yet measured.

An end-to-end worked control cycle

To make the architecture concrete, we walk through a single control tick with the frozen numbers, so the reader can see how the layers of §2, the surrogates of §6, and the clamp of §4 compose in real time. The tick budget is the millisecond tier of figure 22; the protective microsecond loop runs beneath it, asynchronously, and is not gated by any of the steps below.

Step 1 — sense and reconstruct. The diagnostic suite (magnetics, Thomson, interferometry, bolometry, tomography, edge imaging, and the in-winding NV probe) delivers a raw, partially-degraded measurement vector. The graph-network reconstructor 29 refreshes each node from its physical neighbours, imputing any dropped channel; the demonstrated reconstruction AUC is **0.980** (KX-L1-A12), and the fused precursor score is computed in the same forward pass (AUC **0.992**, KX-L1-A20). The output is a calibrated state estimate $\hat{x}^a$ with covariance P^a from the ensemble-Kalman update 33–33, not a point.

Step 2 — predict the shadow. The PINN equilibrium operator returns $\psi(R, Z)$ in **0.87 ms** at **0.139%** RMS- ψ (KX-L1-A11), and the flux surrogate returns the transport-consistent Q and the local fluxes ($R^2 = 0.998$, anchor recovery **3.076** $\rightarrow$ **2.998**, BR-L2-A12). The four twins are advanced one macro-step under the conservation gates 3–4; the shadow now leads the plant by $\tau_{\text{lead}} \geq 50 \text{ ms}$. Because both operators are differentiable, the twin also exposes its Jacobian J_k by automatic differentiation (accepted only when AD matches finite differences to $< 10^{-3}$, SH-032).

Step 3 — decide and gate. The predictive variance and the conformal/OOD monitors 34–36 are evaluated. If the estimate is in-distribution and the variance is below threshold, the learned MPC of 14 computes the nominal command u_{nom} by solving the condensed QP 15 about the current shadow state. If the monitors fire, authority is handed back to the deterministic floor, and the learned layers become advisory — the maturity gate of figure 6 in action.

Step 4 — project onto safety. The candidate u_{nom} passes the three decoupled CBF projections 21. Away from every limit the numerators are zero and $u^* = u_{\text{nom}}$ — the clamp is invisible. Near the β cap the β -channel numerator turns positive and the clamp subtracts exactly enough heating to hold $\dot{h}_\beta = -\gamma h_\beta$, keeping β_N on **3.500** with $h_{\text{min}} = +0.000$ and the command in **[1.06, 1.25]** (the worked example of §4.5). Proposition 4.8 guarantees the state cannot leave $\mathcal{C}$.

Step 5 — actuate, log, learn. The projected command u^* reaches the actuators; the decision, its inputs, the active barriers, and the model versions are written to the signed hash-chained ledger (§9). The measured response feeds the next assimilation step, and, offline, the operator-feedback safe-RL loop ^[10] may refine the policy — but only within the CBF-certified safe set, so exploration never risks the envelope. The cycle repeats.

Two properties of the cycle are worth drawing out. First, it is *predictive*: because the twin leads the plant by $\geq 50 \text{ ms}$, steps 3–4 act on where the plasma will be, so a developing excursion is corrected before it manifests rather than after. Second, it is *honest*: at step 3 the loop can decline to act on a model it does not trust, handing control to the deterministic floor, so the learned layers never command outside the envelope their evidence supports. The combination — act early on a trusted forecast, abstain when the forecast is untrustworthy, and clamp everything to the safe set — is what distinguishes a certifiable autonomous loop from a merely capable one. A conventional high-performance controller can do the first; only the addition of calibrated abstention and a certified clamp gives the second and third, and it is the conjunction of all three that a regulator requires. The cycle also makes explicit where each of the paper’s components lives: the surrogates in step 2, the uncertainty machinery in steps 1 and 3, the performance controller in step 3, the safety projection in step 4, and the provenance infrastructure in step 5 — so the architecture is not an abstract diagram but a concrete sequence executed every control tick.

This five-step cycle is the operational content of the whole paper: the surrogates make it fast enough to run ahead of the plant, the calibrated uncertainty makes it honest about what it knows, and the clamp makes it safe regardless. The following sections deepen the pieces that steps 1–5 invoke: the quantum-limited sensing of step 1 (§8), the verification-first compute that produces every number (§9), and the resource-honest quantum frontier that sits entirely offline (§16).

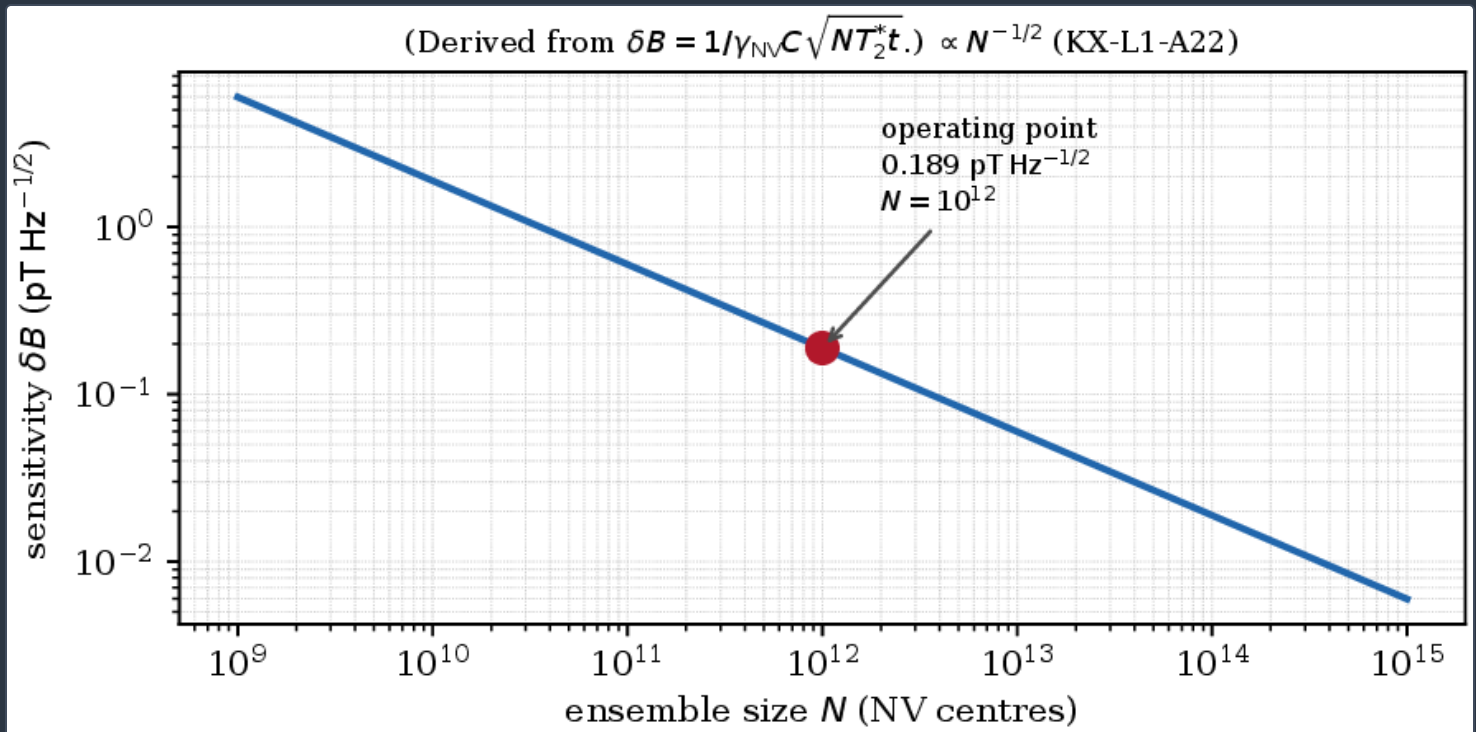
Quantum sensing (Tier 2, near-term)

Magnetometry at or near the quantum limit sharpens the state estimate that feeds the loop. This is the one near-term win in the quantum column, and even here the advantage is bounded — we state where it pays and where it does not. The distinction from Tier 3 is sharp: *quantum sensing* exploits quantum coherence in a measurement device and is a mature, deployable technology today ^[60,61,62]; *quantum computation* exploits quantum coherence in a processor and, for fusion, is an offline calibration tier a decade away (§16). Conflating the two is the most common error in “quantum for fusion” discussion, and keeping them apart is part of the resource-honest posture. The sensing layer enhances the state estimate that Tier-1 control consumes; it never reaches an actuator on its own.

NV-DIAMOND MAGNETOMETRY: THE SHOT-NOISE FLOOR

The in-winding ensemble NV-diamond magnetometer uses Ramsey interferometry on an ensemble of N nitrogen-vacancy centres. The shot-noise-limited sensitivity of such an ensemble with dephasing time T_2^* , interrogation time t , and readout figure of merit C is (Appendix 39)

For the conservative in-winding chip ($N = 10^{12}$, $T_2^* = 1 \mu s$, $C = 0.03$), equation 37 gives $\delta B = 0.189$ pT Hz^{-1/2} at a bandwidth $1/(2\tau_{\text{shot}}) = 125$ Hertz *set by a* $\tau_{\text{shot}} = 4 \mu s$ *Ramsey shot (gate KX – L1 – A22)* *degen2017, barry2020, rondin2014. The* $N^{-1/2}$ *scaling and the operating point are shown in figure 33. The Ramsey sequence is the reason the sensitivity takes this form: a/2* *pulse prepares a superposition of the NV ground – states in sublevels, the spin accumulates a phase = $\gamma_{\text{NV}} t$ over the free – evolution time t , and a second/2 pulse converts that phase to a population readout optically. The optimal* *is set by the ensemble dephasing time T_2^** — — *integrate longer and the coherence is lost, shorter and the phase is under – resolved — — which is why T_2^* appears in 37* and why the picotesla floor is a coherence-limited quantity, not a photon-counting one (the derivation is in Appendix 39).



(Derived from 37.) NV-ensemble shot-noise sensitivity $\delta B \propto N^{-1/2}$; the operating point $\delta B = 0.189$ pT Hz^{-1/2} at $N = 10^{12}$, $T_2^* = 1 \mu s$, $C = 0.03$ (KX-L1-A22).

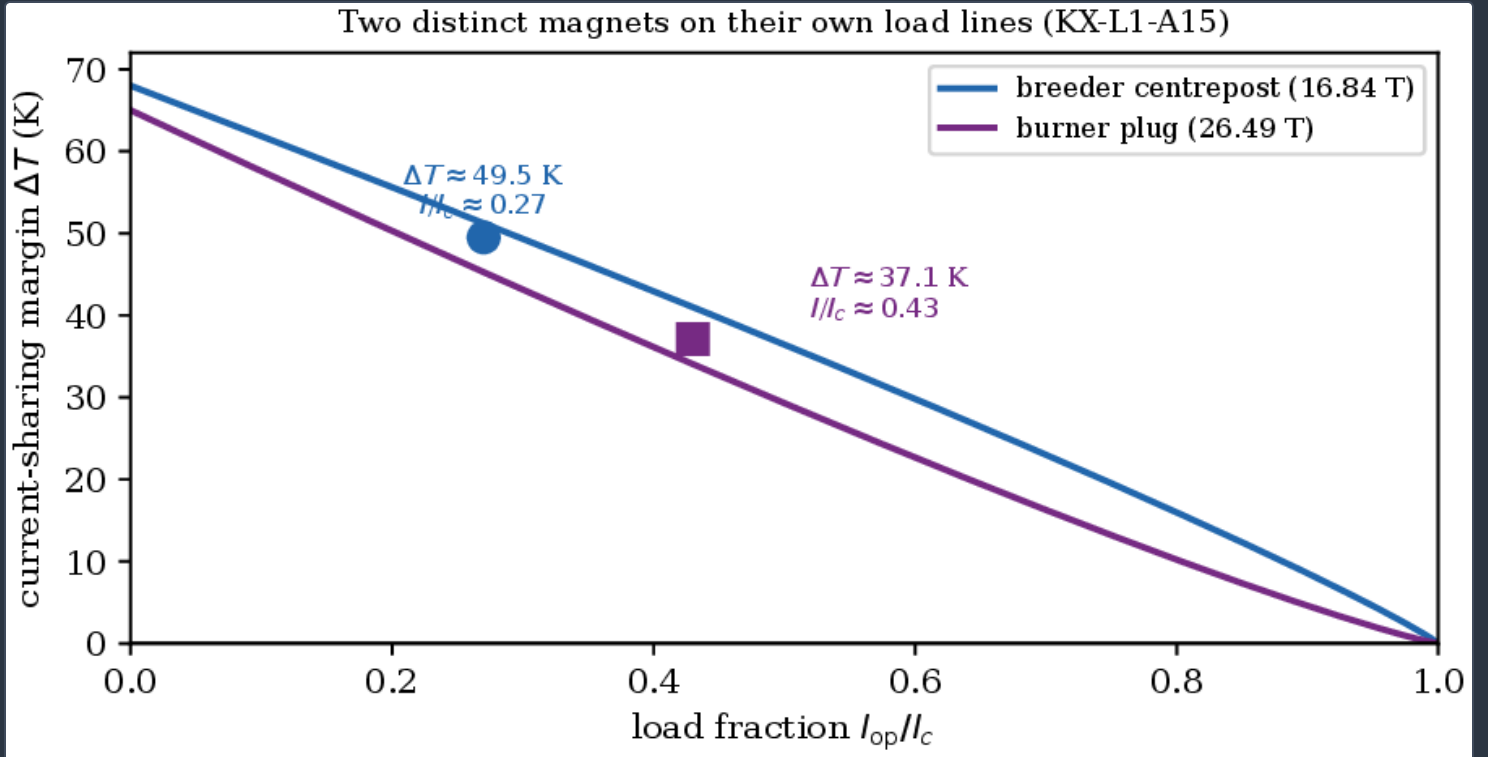
pT Hz^{-1/2} at $N = 10^{12}$,

THE QUENCH PRECURSOR IS SINGLE-SHOT DETECTABLE

A quench that sheds the screening current of a no-insulation turn ($I_{\text{turn}} = 1 \text{ kA}$) produces a local field step

$$\Delta B = \frac{\mu_0 I_{\text{turn}}}{2\pi d} \approx 100 \text{ mT} \quad (d = 2 \text{ mm}),$$

roughly 10^9 times the per-shot floor, so the precursor is single-shot detectable with a per-shot signal-to-noise ratio $\sim 1.5 \times 10^9$ and margin to spare for multiplexing and in-service degradation. The picotesla sensitivity is not *needed* to see a millitesla step; its value is the enormous margin it buys, which is exactly why the protective function survives the $39\times$ NV degradation of SH-037. This is a design principle worth stating generally: sensitivity beyond what the signal requires is not waste but reserve, and in a radiation environment where sensors degrade over their service life, the reserve is what keeps the protective function above its acceptance floor as the sensor ages. A sensor sized exactly to the nominal signal would fail its acceptance test after modest degradation; one sized with 10^9 margin tolerates degradation, multiplexing, and in-service drift simultaneously and still clears the bar. The two magnets are kept distinct throughout: figure 20 plots the current-sharing temperature margin along each load line — the breeder centrepost holds $\Delta T \approx 49.5 \text{ K}$ at load fraction ≈ 0.27 (16.84 T), and the burner plug $\Delta T \approx 37.1 \text{ K}$ at load fraction ≈ 0.43 (26.49 T).



Current-sharing temperature margin along each magnet load line (KX-L1-A15). The breeder centrepost (16.84 T , $\Delta T \approx 49.5 \text{ K}$ at $I/I_c \approx 0.27$) and the burner plug (26.49 T , $\Delta T \approx 37.1 \text{ K}$ at $I/I_c \approx 0.43$) are modelled as the two distinct devices they are.

MULTIPLEXING MARGIN AND RADIATION HARDNESS

The 10^9 per-shot margin is not idle: it is the budget the design spends on multiplexing and on in-service degradation. Spatial multiplexing across the winding trades sensitivity for coverage — reading M locations from one ensemble divides the per-location integration time and raises δB by $\sqrt{M}$ (37) — so even a 10^3 -fold multiplex leaves $\sim 10^6$ SNR on the millitesla step. Radiation hardness is the second draw on the margin: NV centres are known to be robust in a neutron/gamma field, but their contrast C and coherence T_2^* degrade with fluence, raising δB ; the SH-037 result — precursor AUC ≥ 0.95 with the NV channel degraded by $39\times$ — is the model-fusion floor that quantifies how much of the margin can be spent to degradation before the protective function is compromised. The Allan-variance floor sets the long-integration limit: below the drift knee the sensitivity improves as $1/\sqrt{t}$ (white shot noise), above it $1/f$ drift dominates, so the precursor detection is deliberately a short-integration ($4 \mu\text{s}$) measurement that sits in the shot-noise-limited regime and never integrates into the drift floor.

SENSOR-CLASS ASSIGNMENT AND THE ENTANGLEMENT KNEE

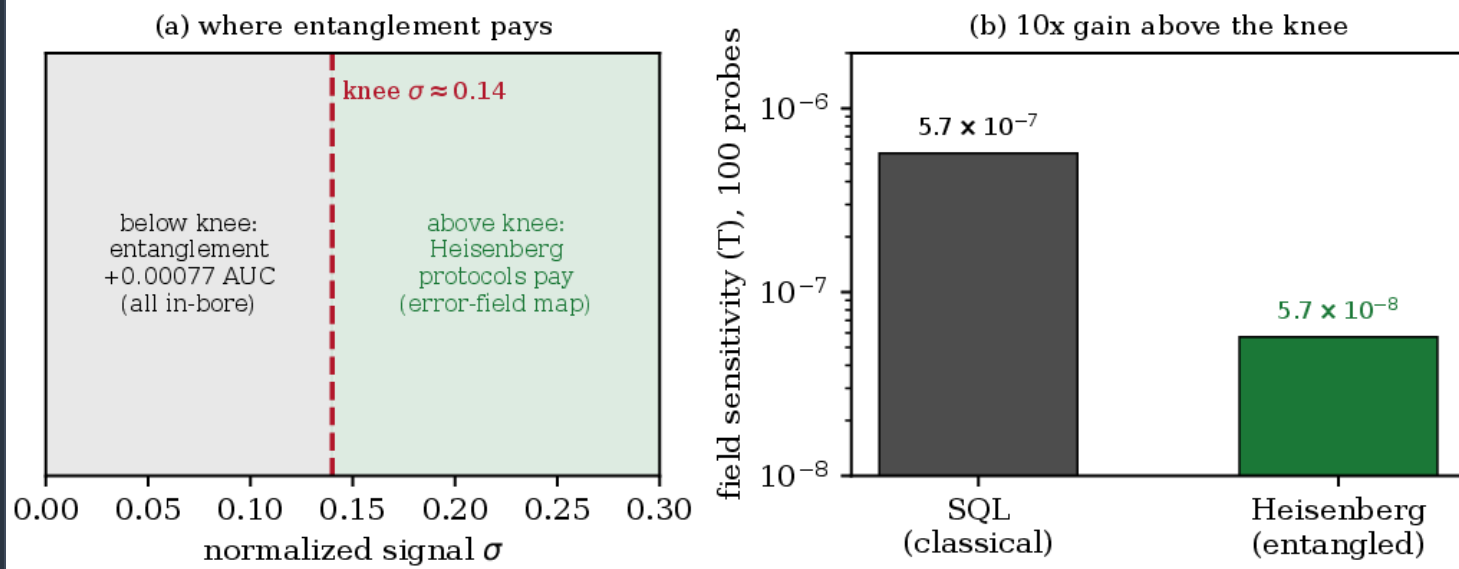
Table 10 assigns sensor classes to the centrepost monitoring requirements. The decisive quantity is where entanglement-enhanced protocols begin to pay. We locate that knee at a normalized signal $\sigma \approx 0.14$ (gate BR-SX-08, figure 21): the in-bore precursor and disruption-mode signals sit *below* the knee, where a coil-plus-SQUID or conventional readout already suffices and entanglement adds only $+0.00077$ AUC — a near-null, kept in the record. *Above* the knee, in deep sub-threshold error-field mapping ($\sigma \approx 0.2$), an NV ensemble running a Heisenberg-limited protocol carries the standard-quantum-limit sensitivity from $5.7 \times 10^{-7} \text{ T}$ down to $5.7 \times 10^{-8} \text{ T}$ at 100 probes and $1 \mu\text{s}$ interrogation. That is the one place quantum sensing earns its complexity, and the chain is specified to use it exactly there and nowhere else ^[19,63]. The knee itself has a simple origin: entanglement (a Heisenberg-limited protocol) improves the sensitivity scaling from $N^{-1/2}$ to N^{-1} in the number of probes, but it also makes the measurement more fragile to decoherence, so the net gain is positive only when the signal is weak enough that the sensitivity, not the decoherence, is the binding constraint. Below the knee the signal is comfortably above the classical floor and entanglement buys almost nothing (the measured $+0.00077$ AUC); above it the task is genuinely sensitivity-limited and the Heisenberg scaling pays. Locating the knee at $\sigma \approx 0.14$ is therefore the decision that tells the sensor chain where to deploy the expensive protocol and where the cheap one suffices.

The sensor-class assignment of table 10 follows a clear logic. The *absolute* field must be known to ≤ 1 ppm for the equilibrium reconstruction to be trustworthy, and only an NMR/EPR probe reaches that accuracy, so it is the primary absolute monitor with $\geq 100\times$ margin on the 1.68 mT task at 10^{-4} relative. A cryogenic Hall probe provides engineering redundancy at coarser accuracy. The quench precursor is an AC signal at the winding, for which a SQUID-plus-pickup channel at femtotesla sensitivity and the NV-diamond ensemble at picotesla are complementary: the SQUID reads at standoff, the NV reads locally and is radiation-hard and cryo-compatible, and the two do not share a failure mode. Atomic optically-pumped magnetometers are reserved for external error-field mapping, above the knee, where their Heisenberg-limited scaling pays. The assignment is thus not “use the most sensitive sensor everywhere” but “match each sensor class to the requirement it uniquely meets” — which is why the picotesla NV is deployed for its *margin* (multiplexing and radiation tolerance) rather than for a sensitivity the millitesla quench step does not demand.

SENSOR CLASS	CLASS-TYPICAL SENSITIVITY	ROLE AT THE CENTREPOST
NMR/EPR probe	0.01–1 ppm absolute	primary absolute field monitor ($\geq 100\times$ margin)
cryogenic Hall probe	$\sim 100 \text{ ppm}$ – 0.1%	engineering redundancy
SQUID + pickup coil	$\sim 1 \text{ fT Hz}^{-1/2}$	AC quench-precursor channel at standoff
NV-diamond ensemble	$\sim 1 \text{ pT Hz}^{-1/2}$	quench / stray-field mapping; rad-hard, cryo
atomic OPM (SERF/scalar)	0.5–20 fT Hz ^{-1/2}	external error-field mapping only

Sensor-class assignment for centrepost field monitoring and quench precursors (class sensitivities from the literature; KX-L1-A18/A22).

Entanglement-enhanced sensing: bounded, above-knee only (BR-SX-08)



Entanglement-enhanced sensing (BR-SX-08). **(a)** the advantage knee at $\sigma \approx 0.14$: in-bore diagnostics sit below it (+0.00077 AUC, near-null); error-field mapping sits above it. **(b)** above the knee a Heisenberg-limited protocol takes the sensitivity from 5.7×10^{-7} to 5.7×10^{-8} T at 100 probes.

Computational methodology: the NVIDIA GPU HPC campaign

The programme is deliberately verification-first, and its compute is described here in full — codes, acceleration, scale and reproducibility — without recourse to any economic framing.

REFERENCE CODES AND THE DE-RISKING CAMPAIGN

Every headline number originates from a frozen reference code, not from the surrogate that later emulates it. The physics backbone spans nonlinear gyrokinetics (CGYRO ^[51], GENE ^[53]), extended-MHD stability, edge/divertor transport, Monte-Carlo neutronics and depletion (OpenMC ^[55]), fast-ion and alpha confinement, quasilinear transport (with neural-network acceleration of the QuaLiKiz class ^[46,47,48]), free-boundary equilibrium (FreeGS/FreeGSNKE ^[54]), and first-principles materials (DFT with machine-learned interatomic potentials). The de-risking campaign ran a large catalog of jobs to closure and stamped each as a frozen anchor in the register ^[1]; the surrogates of §6 are trained *against* these codes and admitted only when they clear the acceptance inequalities. Independent re-run and byte-for-byte provenance are the subject of the verification companion ^[20], and the extended methodology — verification versus validation, the two-tier reproduction gate, and provenance as a safety property — is set out in §10. Table 11 lists the backbone codes and the surrogate each anchors.

The problem scale spans from the tractable to the frontier. Nonlinear gyrokinetic runs are the most expensive element — resolving the turbulent spectrum in five phase-space dimensions over many correlation times — and are run offline on the GPU fleet; the linear spectrum is cheaper and byte-reproduces its reference ($\gamma = +0.318$, BR-L2-A18). Free-boundary equilibria on the 129×129 grid, Monte-Carlo neutronics to statistical convergence, and DFT-with-machine-learned-potentials materials screens complete the backbone. Each code carries its own verification evidence — convergence under grid refinement, reproduction of analytic limits, and, where applicable, comparison to published experiment — before any surrogate is trained against it. The surrogates then compress these expensive calculations into control-latency operators, but only after clearing their acceptance inequalities against held-out cases from the very codes they emulate; the compression is what makes the twin predictive, and the acceptance gate is what makes the compression trustworthy. It is worth stating the ratio plainly: a nonlinear gyrokinetic evaluation is far too slow for a control loop, and even a reduced quasilinear model is a heavy call at kHz rates; the learned operators reproduce the transport response three-to-five orders of magnitude faster than the codes they emulate ^[46,47], and the equilibrium operator runs $2877\times$ faster than the free-boundary Picard solve (KX-L1-A11). That speed-up is what converts an offline design tool into a real-time predictive twin — and it is bought without sacrificing fidelity only because the acceptance inequalities bound the compression error and the abstention monitor forbids extrapolation.

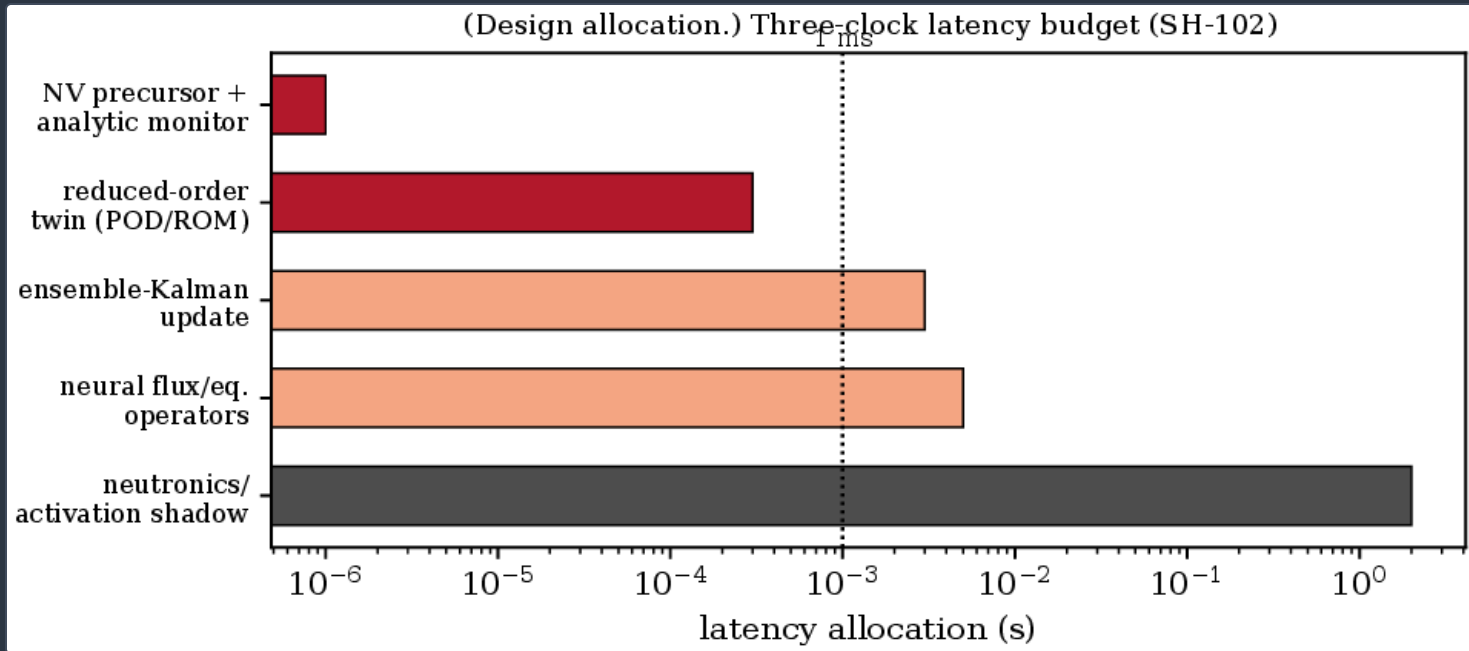
PHYSICS	REFERENCE CODE	SURROGATE / USE
nonlinear gyrokinetics	CGYRO ^[51] , GENE ^[53]	neural flux operator (SH-026)
quasilinear transport	QuaLiKiz class ^[46,47]	degree-3 flux surrogate (BR-L2-A12)
free-boundary equilibrium	FreeGS/FreeGSNKE ^[54]	PINN equilibrium (KX-L1-A11)
Monte-Carlo neutronics	OpenMC ^[55]	neutronics shadow twin
integrated transport	TGYRO	confinement anchor (BR-L1-A11)
first-principles materials	DFT + ML interatomic potentials	alloy screening; ADAPT-VQE check

Reference codes of the de-risking campaign and the learned surrogate each anchors. Every headline number originates from the code, not the surrogate that later emulates it.

GPU ACCELERATION AND THE THREE-CLOCK PARTITION

Training and large offline Monte-Carlo run on the NVIDIA GPU HPC campaign (A100/H100-class silicon); the online digital twin and model inference run on local GPU nodes at the millisecond tier, co-located with the plant for latency; and the hard real-time loop — the compute that reaches an actuator — runs on deterministic bare-metal and FPGA nodes with no hypervisor and no garbage collector. The three tiers are physically separated by role and by security domain and joined only by a physical data diode, so the microsecond protective loop never contends for bandwidth with a model download or a training job. A deterministic, time-sensitive network with sub-microsecond time distribution (PTP/White-Rabbit) gives every node and sensor a common timebase — the precondition for the latency-split clamp of 20 and for lineage timestamps to mean anything. Figure 22 gives the design latency budget by stage: the microsecond NV precursor and analytic monitors are carried by the fast loop, the reduced-order twin at sub-millisecond, the ensemble-Kalman and neural operators at milliseconds, and the neutronics/activation shadow at seconds.

The three-clock partition is a security architecture as much as a latency one. The training tier holds the GPU fleet and the largest attack surface; it is where models are learned but where no actuator is reachable. The online tier holds the digital twin and inference; it can advise and, when a model has cleared its maturity gate, command — but always through the clamp. The protective tier holds the deterministic loop; it has the smallest, most auditable codebase, no network stack, and the sole direct path to the magnet dump. Between the tiers sits a physical data diode — a one-way optical link — so information flows from protective to online to training for logging and learning, but no command flows back down into the protective tier: a compromised training job cannot reach through the diode to delay or spoof a protective action. This is the concrete meaning of “loss of the AI degrades the plant to a safe state”: the protective tier does not depend on the tiers above it, so their failure — whether a crash, an overload, or an intrusion — leaves the deterministic floor intact. The design is a form of defence in depth: an adversary or a fault that reaches the training tier finds no path to an actuator; one that reaches the online tier is still bounded by the clamp and the maturity gate; and one that somehow reached the protective tier would face a codebase small enough to be exhaustively verified and static enough to be formally checked. Each tier’s trust budget is matched to its consequence, and the diode ensures the budgets compose rather than leak. This is why the security architecture and the safety architecture are the same architecture: both follow from the principle that the fastest, most consequential action must depend on the least, most-verifiable machinery.



(Design allocation.) The three-clock latency budget by stage (SH-102): microsecond protective loop, sub-millisecond reduced-order twin, millisecond assimilation and neural operators, second-scale neutronics shadow. The fast protective loop is never gated behind the slower learned layers.

REPRODUCIBILITY, VERIFICATION, AND MLOPS

The pipeline is built for byte-level reproducibility (figure 23). Every model is versioned with its data and configuration and promoted only through a validation gate — Tier-1 byte-exact reproduction and Tier-2 tolerance reproduction against an independent library — and a model registry tracks what is deployed where, at which version, with which validator ^[14]. Every training example carries its provenance (which pulse, which simulation, which device), and provenance-gated calibration is load-bearing for safety, not paperwork: control-barrier invariance held in **10/10** provenance-gated trials against **1/10** ungated. Every runtime decision is written to a signed, hash-chained ledger ^[8]. No result in this paper depends on unpublished or proprietary data, and each is regenerable from the archived scripts and the register.

The hash-chained ledger is what makes the audit trail tamper-evident. Each entry contains the decision, its inputs, the active barriers, the model versions, and the cryptographic hash of the previous entry, so any alteration of a past entry invalidates every subsequent hash — an auditor can verify the chain end to end and detect any post-hoc modification. Combined with the provenance stamps on the training data and the two-tier reproduction gate, this closes the loop from a training datum, through a validated model version, to an actuation, to an immutable log of that actuation. This is the concrete infrastructure behind the claim that the system is *auditable*: not that it produces documentation, but that every runtime decision is cryptographically bound to the verified artifacts that produced it — the difference, for a regulated plant, between a safety case that describes the intended system and one that can be checked against the running system. Table 12 states the model-promotion stages.

STAGE	GATE	ARTIFACT RECORDED
train	held-out error within bound	data hash, config, weights
verify (Tier-1)	byte-exact reproduction of the anchor	reproduction log
validate (Tier-2)	tolerance match vs independent library	comparison report
register	provenance complete	registry entry (version, validator)
deploy	maturity-gate rung cleared	signed-ledger genesis record
monitor	runtime gates hold (OOD, coverage, variance)	per-decision ledger records

Model-promotion stages in the MLOps pipeline. A model advances only by clearing the gate at each stage; the runtime authority it may then hold is capped by the maturity gate (table 4).

(Schematic.) The verification-first assurance stack*verification-first: nothing acts before it clears its gate*

frozen reference codes (CGYRO, OpenMC, FreeGS, DFT+MLIP)

de-risking gates: independent re-run, stamped in register

surrogates trained against codes; acceptance inequalities

gated model promotion (byte-exact / tolerance) + registry

runtime: signed hash-chained ledger + maturity gate

(Schematic.) The verification-first assurance stack: frozen reference codes → independently re-run de-risking gates → surrogates trained against the codes and admitted by acceptance inequalities → gated model promotion and registry → signed-ledger runtime with the maturity gate.

Extended verification and the de-risking methodology

The claim that every headline number is a frozen, reproducible anchor is only as strong as the methodology behind it. This section states that methodology, because it is what distinguishes a de-risked design point from a plausible one.

VERIFICATION, VALIDATION, AND THE TWO-TIER REPRODUCTION GATE

We use the terms in their precise sense. *Verification* asks whether the code solves the equations correctly — a mathematics question, answered by convergence studies, method-of-manufactured-solutions tests, and reproduction of analytic limits. *Validation* asks whether the equations describe the machine — a physics question, answered against experiment where available and against higher-fidelity codes otherwise. The de-risking campaign runs both and then adds a third, sharper requirement: *independent reproduction*. Each gate is re-run from its archived inputs and must reproduce its frozen result under one of two tiers — Tier-1 byte-exact reproduction, where the recomputation matches the stored result bit-for-bit, or Tier-2 tolerance reproduction, where an *independent* library reproduces the result within a stated tolerance. A gate that clears neither is not frozen. The control-relevant example is the CGYRO linear spectrum: it byte-reproduces its reference growth rate $\gamma = +0.318$ (BR-L2-A18), a Tier-1 pass, which is why the flux surrogate can be anchored to it with confidence.

PROVENANCE AS A SAFETY PROPERTY

Provenance is not bookkeeping; it is load-bearing for safety. Every training example carries its origin — which pulse, which simulation, which device, which code version — and every model is promoted only through the reproduction gate, so a model in the loop is guaranteed to be the validated one. The campaign measured the cost of skipping this: control-barrier invariance held in **10/10** provenance-gated trials against **1/10** ungated, the difference being that an ungated model may silently be a stale or mis-configured checkpoint whose error bound no longer holds. Because the CBF guarantee (Proposition 4.8) assumes a bounded surrogate error, a model whose provenance is unknown is a model whose error bound is unknown, and the guarantee is void for it. Provenance gating is therefore the precondition that makes the safety theorem apply to the deployed system, not merely to its idealisation. Every runtime decision is written to a signed, hash-chained ledger (Appendix 50), so the chain from a training datum through a model version to an actuation is tamper-evident end to end.

WHAT “FROZEN” MEANS FOR THIS PAPER

Every quantitative claim in this paper is a frozen anchor in the register ^[1] with a serial identifier given inline, and every figure plots only reported summary quantities or analytic curves derived from stated equations — no figure draws invented data as a measured result, and schematic figures are captioned as such. The figure-generation scripts are archived with the deposit, so each figure is regenerable from the frozen numbers. This is the sense in which the paper is a set of reproducible results rather than assertions: a reader with the register and the scripts can regenerate every number and every figure. The verification-first posture is the reason the quantum negatives are reported as negatives — they are frozen results too, and a resource-honest programme freezes its misses with the same discipline as its hits.

Diagnostics and real-time state estimation

The control loop begins with the state estimate, and in a **14 MeV** neutron field the estimate must survive the loss of individual diagnostics. This section deepens the reconstruction layer that step 1 of the worked cycle (§7) invokes.

THE DIAGNOSTIC GRAPH

The diagnostic suite is organised as a graph whose nodes are analyzers (magnetics, Thomson scattering, interferometry with fringe-jump recovery, bolometry, tomographic inversion, edge imaging, and the in-winding NV probe) and whose edges connect analyzers that constrain overlapping physical quantities (table 13). Because neighbouring channels carry redundant information about the same equilibrium, the message-passing update 29 reconstructs a dropped node from its neighbours rather than losing it; the demonstrated reconstruction AUC is **0.980** at NRMSE **0.396** on synthetic data (KX-L1-A12) ^[12]. The graph is not a convenience: it is the structural encoding of “no single sensor is load-bearing,” which is why the estimate degrades gracefully rather than catastrophically when a channel fails. The choice of edges encodes physics: two analyzers are connected when they constrain a common quantity, so information propagates along physical-meaningful paths rather than through a dense all-to-all coupling. This restricts the reconstruction to the manifold of physically-plausible states, which improves sample efficiency and makes each reconstructed value traceable to the neighbours that supplied it — an interpretability property a black-box imputer lacks and a regulator requires. The reconstruction head and the precursor head share the graph encoder, so the state estimate and the off-normal score are produced in one forward pass, keeping the whole reconstruction inside the millisecond budget.

node	constrains	metric / status	
magnetics (drift-corrected)	boundary field, I_p , position	stable over pulse	
Thomson scattering	T_e, n_e profiles	profile fit within uncertainty	
interferometry (fringe recovery)	line-integrated n_e	robust to fringe jumps	
bolometry / tomography	emissivity, radiated power	inversion within tolerance	
edge imaging	strike point, detachment state	tracks divertor state	
NV-diamond (in-winding)	local field, quench precursor	0.189	pTHz^{-1/2} (KX-L1-A22)
GNN reconstructor	full state + precursor	AUC 0.980 (KX-L1-A12)	

The diagnostic suite as graph nodes, the quantities each constrains, and its acceptance metric or status (companion ^[12], serials in the register).

VIRTUAL-SENSOR MESH AND DROPOUT ROBUSTNESS

Two mechanisms give dropout robustness. First, *redundancy across modalities*: optical, magnetic, and quantum channels fail in uncorrelated ways, so the loss of any one leaves the others to constrain the state. Second, a *dropout-robust filter* whose admission to authority is gated by a seeded-dropout test — it must hold reconstruction error within tolerance and keep state error bounded with N channels removed before it acts (SH-042/043). The maturity gate of §2 caps authority until that bar is met, the same principle that governs the certified clamp. The end-to-end consequence is the SH-037 result: the protective precursor retains $AUC \geq 0.95$ even with the NV channel degraded by **39×** (figure 17). The virtual-sensor mesh is the concrete realisation of the second mechanism: for each physical diagnostic the reconstructor learns a virtual estimate from the rest of the graph, so a dropped channel is replaced by its virtual counterpart in real time rather than leaving a hole in the state vector. The mesh is validated by a leave-one-out test — ablating each channel in turn and checking that the reconstructed value tracks the withheld measurement within tolerance — and because it is trained with seeded dropout (Appendix 47) it degrades gracefully as channels are lost rather than failing at a threshold.

PRECURSOR LIBRARY AND CROSS-MACHINE TRANSFER

The off-normal precursor score is trained against a curated precursor library with leakage-free splits, and a disruption model pre-trained on TCV, DIII-D and NSTX transfers to the compact-ST breeder with honest transfer uncertainty^[12]. Cross-machine transfer is the reason a first-of-a-kind machine can carry a credible disruption predictor from its first shots rather than accumulating a training set from scratch; the transfer uncertainty is propagated into the same conformal/OOD gate 34–36 that governs the surrogates, so an out-of-distribution transfer throttles its own authority. Per-decision explainability — a do-calculus attribution of which signals drove an alarm — is a regulator-facing requirement and is reported in the causal-explainability companion.

Cross-machine transfer deserves emphasis because it addresses the first-of-a-kind problem directly. A disruption or precursor model trained only on the target machine cannot exist before the target machine runs; but the physics of the precursor — the current-redistribution signature, the mode structure, the approach to a limit — is shared across devices, so a model pre-trained on TCV, DIII-D and NSTX carries usable structure to the compact-ST breeder from its first shots. The transfer is not assumed to be perfect: the domain shift between a conventional aspect-ratio machine and $A = 2.5$ is real, and it is quantified as an added epistemic-uncertainty term that inflates the predictive variance in the gate 34. So the transferred model is admitted only at the authority its *transfer uncertainty* warrants, and it earns higher authority as target-machine data accumulate and the transfer uncertainty shrinks — the same maturity-gated principle applied to the first-of-a-kind case. This is how a new machine can carry a credible protective model on day one without over-trusting it. The alternative — waiting to accumulate a target-machine training set before deploying any protective model — is untenable for a nuclear plant, which needs disruption protection from its first high-current shots. Transfer learning with honest transfer uncertainty resolves the tension: the machine is protected from day one by a model that knows how much to trust itself, and the protection improves as data accumulate and the transfer uncertainty shrinks, all within the same maturity-gate discipline that governs every other learned component.

The magnet-protection subsystem in depth

The magnet-protection function is the demonstrated exemplar of the whole pattern — a fast surrogate, a fused multi-sensor estimate that beats any single channel, and an early precursor that fires before the conventional signal — and it is where the Tier-1 and Tier-2 layers meet. We derive its detection statistic and its operating points here.

THE SLOW-PROPAGATION QUENCH PROBLEM

The Hyperion magnets are wound from no-insulation (NI) REBCO coated conductor, the technology behind record fields ^[28] and compact high-field tokamak concepts ^[23,22]. NI REBCO has a large temperature margin and a very slow normal-zone propagation velocity (NZPV), of order millimetres per second. This is a protection liability: a developing hot spot can reach a damaging temperature before it generates a measurable resistive voltage, so conventional voltage-tap detection is ineffective ^[64]. The subsystem therefore fuses non-voltage observables — a fibre-Bragg strain-rate channel that reads the mechanical signature of local current redistribution, and a magnetization channel that reads the collapse of the screening current — and adds the in-winding NV magnetometer of §8 for a local, vector precursor.

Quantitatively, the danger is a race between two timescales. The hot-spot temperature rises on the adiabatic timescale set by the local Joule heating and the enthalpy of the conductor, while the resistive voltage a voltage-tap detector would see grows only as the normal zone spreads at the NZPV of order millimetres per second. In a large-margin NI conductor the temperature can reach a damaging value before the normal zone is large enough to produce a detectable terminal voltage — the detection lag is the liability. The non-voltage channels break the race because they see the *local* precursor on the millisecond timescale of the developing hot spot: the strain-rate channel reads the mechanical signature of the local current redistribution as it begins, the magnetization channel reads the collapse of the local screening current, and the NV magnetometer adds a direct local field measurement. Fusing them is what turns a set of weak, local, early signals into a trip decision with the demonstrated low-false-positive performance.

QUADRATURE FUSION AND THE ROC

Each channel produces a matched-template z -score over the precursor window: a strain-rate score z_S and a magnetization score z_M , each with approximately unit-variance noise and signal-dependent mean (window-statistic SNR). Because the two modalities fail in uncorrelated ways, their quadrature fusion $z_F = (z_S + z_M)/\sqrt{2}$, has noise variance $\frac{1}{2}(\text{Var } z_S + \text{Var } z_M) = 1$ and mean $(\mu_S + \mu_M)/\sqrt{2}$, so the fused window-statistic SNR is

$$\text{SNR}_F = \frac{\text{SNR}_S + \text{SNR}_M}{\sqrt{2}}.$$

On the frozen 2×10^4 -trial-per-class synthetic study in a deliberately weak-fault regime (per-sample SNR set $10\text{--}30\times$ below the design-basis fault, where every detector saturates), the single-channel window SNRs are $\text{SNR}_S = 2.00$ and $\text{SNR}_M = 2.78$, so 39 gives $\text{SNR}_F = (2.00 + 2.78)/\sqrt{2} = 3.38$ (KX-L1-A20). For a Gaussian shift model the area under the ROC of a detector with window SNR d' is

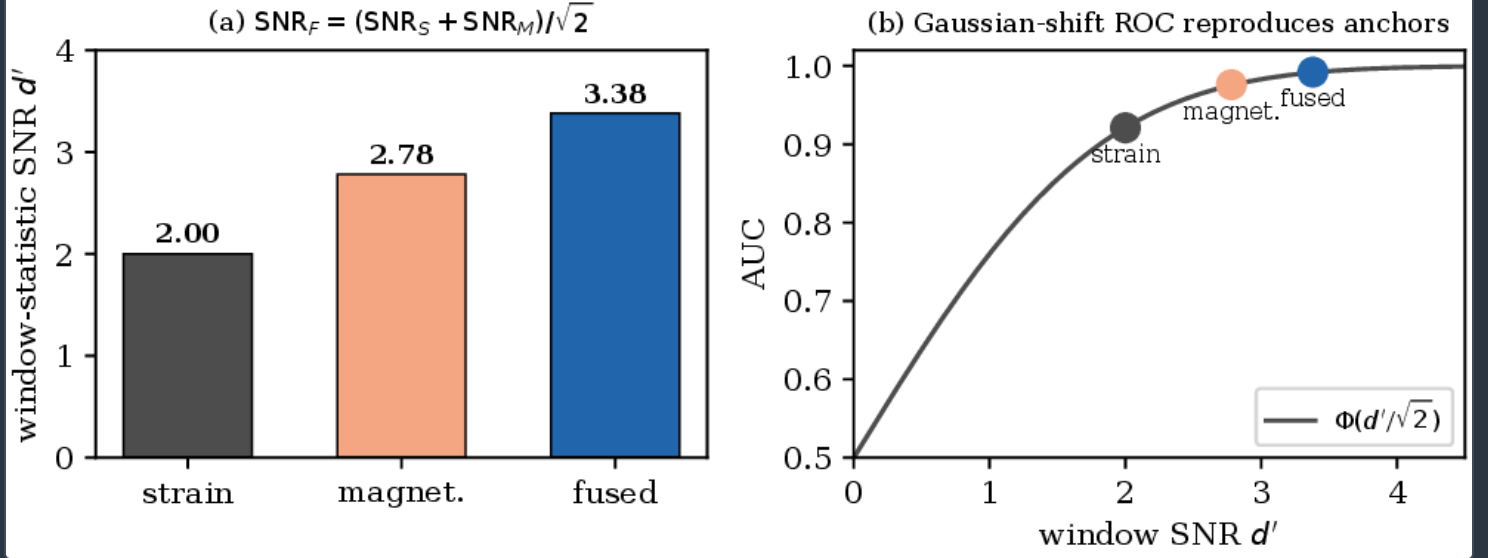
$$\text{AUC} = \Phi(d'/\sqrt{2}),$$

with Φ the standard normal CDF, which reproduces the frozen anchors exactly: $\Phi(2.00/\sqrt{2}) = 0.921$ (strain), $\Phi(2.78/\sqrt{2}) = 0.976$ (magnetization), $\Phi(3.38/\sqrt{2}) = 0.992$ (fused) (figure 24). The fusion gain is largest where a trip decision lives — the low-false-positive corner: at 1% false-positive rate the fused true-positive rate is 0.857 against 0.671 and 0.372 for the single channels, and at 0.1% it is 0.597 against 0.386 and 0.135 (table 14, figure 17). The different channels fail in different ways, and their fusion is what buys the early, low-false-alarm warning.

CHANNEL	AUC	TPR @ 1% FPR	TPR @ 0.1% FPR
fibre-Bragg strain-rate	0.921	0.372	0.135
magnetization	0.976	0.671	0.386
fused (strain-rate + magnetization)	0.992	0.857	0.597

Fused quench-precursor detection (KX-L1-A20; synthetic, 2×10^4 trials/class). Window-statistic SNRs 2.00/2.78/3.38 follow the quadrature rule 39; AUCs follow 40. The fused detector dominates the full ROC.

Quadrature fusion of quench precursors (KX-L1-A20)



Quadrature fusion of quench precursors (KX-L1-A20). (a) the fused window SNR follows the quadrature rule 39, $3.38 = (2.00 + 2.78)/\sqrt{2}$. (b) the Gaussian-shift ROC relation $AUC = \Phi(d'/\sqrt{2})$ 40 reproduces the three frozen anchors exactly. Derived from the frozen window SNRs; no ROC data are invented.

CURRENT-SHARING AND THE PROTECTIVE OBSERVABLE

The protective observable is the current-sharing temperature margin. In a REBCO conductor carrying operating current I_{op} at field B , the critical current $I_c(T, B)$ falls with temperature; the current-sharing temperature T_{cs} is where $I_c(T_{cs}, B) = I_{op}$, above which current begins to share into the normal metal and Joule heating starts. The margin $\Delta T = T_{cs} - T_{op}$ is how much headroom the conductor has before that onset. A developing hot spot, a lost cooling channel, or a fluence-degraded I_c all show up as a drift of ΔT toward zero, so the twin's protective signal is the modelled-versus-measured drift of ΔT along the load line. The slow normal-zone propagation of NI REBCO means this drift can develop locally before it produces a global voltage — which is exactly why the non-voltage strain and magnetization channels, and the local NV field measurement, are needed: they see the local precursor that the terminal voltage misses. The fused detector of 30–40 is the statistical machinery that turns those local, noisy precursor signals into a trip decision with the demonstrated AUC **0.992** and the low-false-positive operating points a protection trip demands. The margin analysis follows the critical surface: along the load line $I_c(T, B)$ the current-sharing temperature $T_{cs}(I_{op}, B)$ is read off where the load current equals the critical current, and the twin propagates the two magnets' surfaces separately (figure 20). The centrepost's larger margin ($\Delta T \approx 49.5\text{ K}$ at load fraction ≈ 0.27) reflects its lower field and lower load fraction; the plug's smaller margin ($\Delta T \approx 37.1\text{ K}$ at ≈ 0.43) reflects its higher field and load. Tracking the modelled-versus-measured drift of ΔT against these surfaces is how the twin turns a static margin into a dynamic protective observable.

TWO DISTINCT MAGNETS ON THEIR OWN LOAD LINES

The subsystem treats the breeder centrepost and the burner plug coil as the distinct devices they are (KX-L1-A15). Along each coil load line the current-sharing temperature margin $\Delta T = T_{cs} - T_{op}$ is the protective observable (figure 20): at the design current the centrepost holds $\Delta T \approx 49.5\text{ K}$ at load fraction $I_{op}/I_c \approx 0.27$ and peak field **16.84 T**, while the plug — a distinct, higher-field magnet — holds $\Delta T \approx 37.1\text{ K}$ at load fraction ≈ 0.43 and peak field **26.49 T**. The twin's protective observable is the modelled-versus-measured drift of ΔT toward the current-sharing line; keeping the two magnets separate — never conflating the centrepost field with the plug field — is a load-bearing element of the magnet-integrity case, and the twin models each on its own critical surface.

Neutronics, activation, and the fuel-cycle shadow

The neutronics twin shadows the 14 MeV source and everything downstream — the tritium-breeding balance ($\text{TBR} = 1.42$), the activation inventory, and the shutdown dose — updated as the operating point moves ^[55]. It runs on the slow clock of the latency budget (figure 22): activation build-up and inventory bookkeeping evolve over hours to months, not milliseconds, so the neutronics shadow informs scheduling and maintenance rather than the fast loop.

COUPLING INTO THE SHADOW

The neutronics domain is coupled to the others through 1–1: the equilibrium $\psi(\mathbf{R}, \mathbf{Z})$ from the MHD twin sets the source geometry $\mathcal{S}_n$, and the volumetric nuclear heating q_n''' drives the thermomechanics twin. Because the source depends on the operating point, the breeding ratio, the activation, and the dose all move as the plasma moves, and the twin propagates that dependence at the slow rate while the conservation gates 3–4 keep the coupled shadow self-consistent. The OpenMC-class Monte-Carlo neutronics ^[55] that anchors this twin is the reference code; a reduced surrogate carries it at the shadow rate, admitted under the same acceptance discipline as the transport surrogate.

REAL-TIME TRITIUM ACCOUNTANCY

A tritium-accountancy twin tracks the in-vessel inventory by a mass balance $dI_T/dt = B_T - \dot{N}_{\text{burn}} - \dot{N}_{\text{perm}} - \dot{N}_{\text{decay}}$, with breeding source B_T (set by TBR and the neutron rate), burn rate $\dot{N}_{\text{burn}}$, permeation loss $\dot{N}_{\text{perm}}$ through the first wall, and radioactive decay $\dot{N}_{\text{decay}} = I_T \ln 2/t_{1/2}$. Self-sufficiency requires the net breeding to close the balance with margin, which is why the frozen $\text{TBR} = 1.42$ carries headroom above unity. The accountancy twin is a bookkeeping function, not a control loop, but it feeds the scheduler and is subject to the same provenance and maturity discipline as the fast-loop models.

REMAINING-USEFUL-LIFE PROGNOSTICS FOR THE CONSUMABLE CENTREPOST

The centrepost accumulates neutron fluence and is a consumable, so a remaining-useful-life (RUL) prognostic schedules its replacement. A dual-channel fluence-driven model tracks two damage measures — displacement damage and the degradation of the superconductor's current-sharing margin — and estimates the RUL as the time until either crosses its limit, with an uncertainty band from the propagated fluence uncertainty. The prognostic is a slow-clock function that turns the neutronics shadow into a maintenance schedule, closing the loop from the physics twin to the operations layer. No economics enters any of this bookkeeping; the twin tracks physical inventory, dose, and damage, and any downstream analysis is reported separately.

The human-machine layer

Autonomy does not remove the operator; it changes the operator's role from moment-to-moment actuation to supervision, intervention, and audit. The experience layers L5–L6 of the stack (§2) provide that role with three functions, each governed by the same maturity gate and the same clamp.

NATURAL-LANGUAGE-SUPERVISED CONTROL

An operator copilot accepts natural-language intent — “raise the density fraction toward the cap while holding β_N below 3.4” — and translates it into setpoints for the reference governor, never into direct actuation. The copilot is grounded: it answers with cited references to the register and the twin state, and it is constrained so that it can never bypass the clamp or command outside the maturity-gated authority of the underlying components. Its acceptance test is intent accuracy on scripted scenarios and the invariant that it never issues a command the clamp would reject; it is admitted at the advisory rung until it clears that bar. The value is that an operator can express a goal at the level of physics rather than of coil currents, while the safety and authority machinery beneath is unchanged — the copilot is a convenience layer, not a control layer.

AUTONOMOUS ROOT-CAUSE ANALYSIS

When an off-normal event occurs, the root-cause analyser assembles the relevant ledger records, the active barriers, and the model versions in the loop at the time, and produces a causal account of what drove the event — which signal tripped, which barrier activated, which model was in authority. Because every runtime decision is in the signed ledger (Appendix 66), the analysis is reconstructive from an immutable record rather than speculative, and its per-decision explainability (a do-calculus attribution of which signals drove the outcome) is the regulator-facing artifact. This turns post-event analysis from a forensic exercise into a query against a complete, tamper-evident log.

THE DIGITAL OPERATOR LOGBOOK

The logbook records the pulse history, the interventions, the maturity-rung changes, and the model-version deployments, all cross-linked to the ledger. It is the operational memory of the plant and the substrate for the cross-machine transfer of §11: the accumulated, provenance-stamped operating history is what a future model trains on, and what an auditor reviews. The human-machine layer thus closes the loop from autonomous operation back to human oversight without reintroducing the operator into the microsecond protective path — supervision at the timescale of decisions, autonomy at the timescale of actuation.

This division of labour is the right one for a regulated plant. The operator is neither displaced nor asked to supervise at a timescale no human can meet; instead the operator sets goals, reviews interventions, and audits the record, while the machine handles the microsecond-to-millisecond actuation the operator could never reach. The copilot lowers the barrier to expressing goals correctly, the root-cause analyser lowers the barrier to understanding what happened, and the logbook preserves the institutional memory across shifts and across the machine's life. None of the three touches the safety layer: they inform and record, but the clamp and the maturity gate are unchanged whether the operator is expert or novice, present or delegating to the copilot. That invariance — the safety guarantee does not depend on the human in the loop — is what makes the human-machine layer a convenience rather than a liability.

The burner control problem

The same architecture drives the D–³He tandem-mirror burner, but the plant is different enough that the safe set, the actuators, and the physics models are all re-expressed. Stating the burner case makes concrete the claim that the contribution is a method, not a point solution (§5.8).

A DIFFERENT CONFINEMENT CONCEPT, THE SAME CONTROL SKELETON

A tandem mirror confines the central-cell plasma electrostatically: high-field end plugs raise an ambipolar potential that reflects central-cell ions, so confinement is set by the plug-to-cell potential rather than by a closed magnetic surface. This is a fundamentally open-field-line topology, the opposite of the breeder's closed flux surfaces, and the control problem inherits the difference: where the breeder must avoid disruptions of a confined current, the burner must sustain a potential structure against end losses and microstability. The frozen anchors reflect a genuinely low-neutron D–³He regime — $f_n = 5.44\%$ against the breeder's neutron-rich D–T — which is what makes direct energy conversion at $\eta_{\text{DEC}} = 0.49$ attractive: most of the fusion power emerges as charged particles that a direct converter can capture without a thermal cycle. The burner's frozen anchors are the engineering gain $Q_E = 1.318$, the neutron fraction $f_n = 5.44\%$ (a genuinely low-neutron regime for the D–³He fuel), the plug field 26.49 T , the direct-energy-conversion efficiency $\eta_{\text{DEC}} = 0.49$, a central-cell $\beta_c = 0.55$, and a plug-to-cell density ratio $n_{\text{plug}}/n_{\text{cell}} \approx 16$. These are the burner analogues of the breeder anchors of table 20, and every learned surrogate and every safety barrier is verified against them.

THE BURNER SAFE SET

Where the breeder's safe set is (f_G, β_N, q_{95}) , the burner's is expressed over the plug-potential margin, the central-cell β_c , and the microstability margins (DCLC and Alfvén ion-cyclotron modes). The barriers take the same half-space form as 9–12: a plug-potential barrier $h_\phi = \phi_i - \phi_i^{\min}$ that keeps the ion-confining potential above the value required for confinement, a β_c barrier $h_{\beta_c} = \beta_c^{\max} - \beta_c$ that keeps the central cell below its MHD ceiling, and microstability barriers that keep the distribution inside the DCLC/AIC-stable window. The clamp of §4 acts on this safe set exactly as on the breeder's, with the same closed-form projection 20 and the same forward-invariance guarantee (Proposition 4.8); only the barrier functions and the actuator map differ.

BURNER-SPECIFIC ACTUATORS AND SURROGATES

The burner actuators are the electron-cyclotron heating that sustains the plug thermal barrier, the sloshing-ion injection that shapes the ambipolar potential, and the direct-energy-converter grid bias that dispatches the charged-particle power. Each has a companion control-suite treatment; the point for the umbrella is that the learned-surrogate machinery transfers: a differentiable surrogate of the plug-potential response replaces the equilibrium operator, a mirror-transport surrogate replaces the flux operator, and the same ensemble-Kalman assimilation and conformal gating carry the calibrated uncertainty. The magnet-protection subsystem transfers too, with the plug coil at 26.49 T on its own load line (figure 20) — distinct from the breeder centrepost, and never conflated with it.

QUANTITY	BREEDER	BURNER
gain	$Q = 3.076$	$Q_E = 1.318$
peak magnet	centrepost 16.84 T	plug 26.49 T
safe-set coordinates	(f_G, β_N, q_{95})	$(\phi_i, \beta_c, \text{DCLC/AIC})$
β observable	$\beta_N \leq 3.5$	$\beta_c = 0.55$
neutron / conversion	TBR = 1.42	$f_n = 5.44\%, \eta_{\text{DEC}} = 0.49$
key ratio	—	$n_{\text{plug}}/n_{\text{cell}} \approx 16$

The burner anchors and the control quantities they set, alongside the breeder analogues. The architecture is identical; the plant, actuators, and barriers differ.

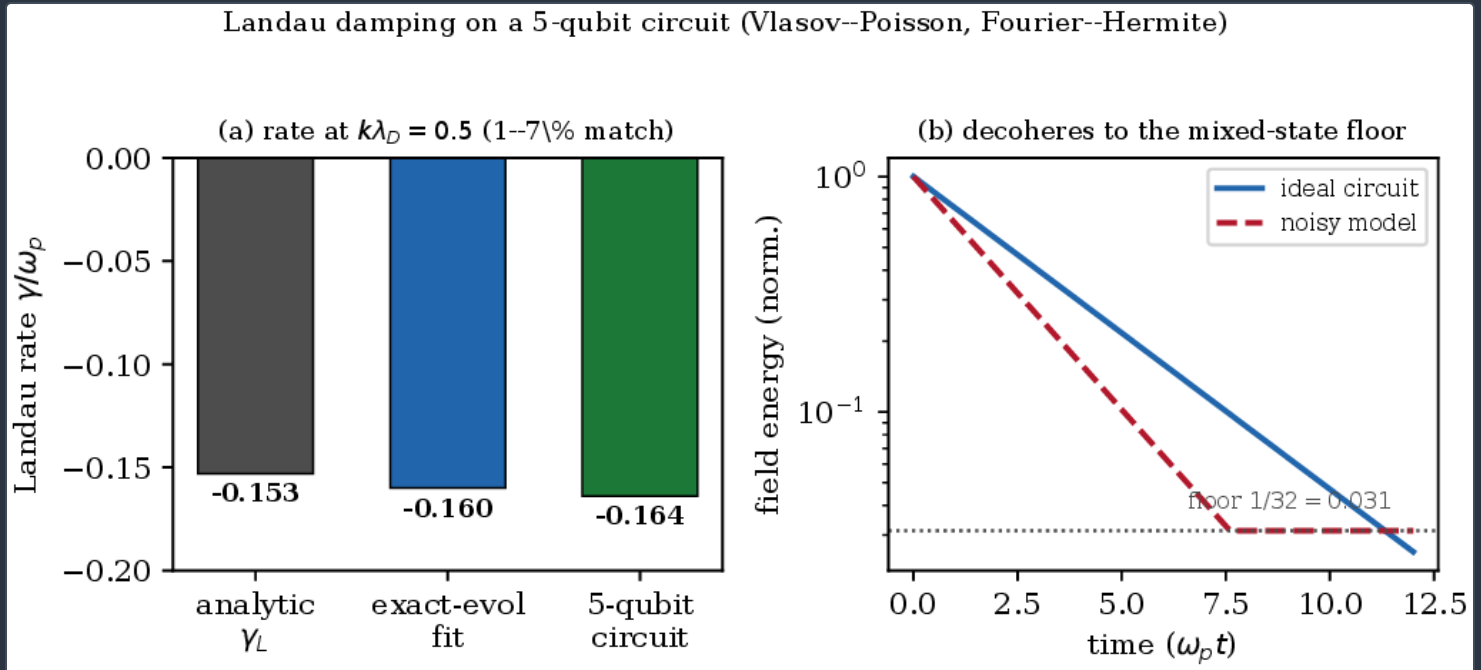
That one certifiable control architecture serves both a closed-field spherical tokamak and an open-field tandem mirror — two confinement concepts whose physics could hardly be more different — is the strongest single piece of evidence that the contribution is a transferable method (§22).

The resource-honest quantum frontier

We place quantum computation where it actually is for fusion, and report the negatives. There is no resource crossover this decade; the classical baseline is cheaper at every demonstrated scale.

QUANTUM SIMULATION OF KINETICS: LANDAU DAMPING ON FIVE QUBITS

The linear kinetic problem maps cleanly to Hamiltonian simulation. Expanding the linearised electrostatic Vlasov–Poisson system (one spatial and one velocity dimension, single Fourier mode, Maxwellian background) in a Fourier–Hermite basis and augmenting it with the electric-field variable yields a Hermitian generator H , so that e^{-iHt} is unitary and is a legitimate Hamiltonian-simulation target (Appendix 35, cf. Engel *et al.* ^[81] and the plasma-quantum reviews ^[82,83]). Evolving this generator on a 5-qubit statevector simulator reproduces the analytic Landau damping rate to 1–7% across three wavenumbers: at $k\lambda_D = 0.5$ the ideal circuit measures $\gamma = -0.164\omega_p$ against the analytic $\gamma_L = -0.153\omega_p$ (ratio 1.07, matching the exact-evolution fit -0.160), and the dispersion benchmark $\omega/\omega_p = 1.4157 - 0.1534i$ reproduces the textbook $1.4156 - 0.1533i$ (figure 25a). Under a realistic noise model (single-qubit error 10^{-3} , two-qubit 10^{-2} , readout 10^{-2}) the field-energy observable decoheres to the maximally-mixed floor $1/32 = 0.031$ and the damping signature is lost (figure 25b): the algorithm is correct, but present-day noisy hardware cannot yet run it ^[17].



Landau damping on a 5-qubit circuit. (a) damping rate at $k\lambda_D = 0.5$: analytic -0.153 , exact-fit -0.160 , circuit -0.164 (1–7% match). (b) the field-energy observable decoheres to the mixed-state floor $1/32$ under a realistic noise model. Frozen anchors (companion ^[17]).

The noise model and why the signature is lost.

The pre-registered hardware negative is not vague: it follows from an explicit noise model. Each single-qubit gate is followed by a depolarising channel $\rho \mapsto (1 - p_1)\rho + \frac{p_1}{3}(X\rho X + Y\rho Y + Z\rho Z)$ with $p_1 = 10^{-3}$, each two-qubit gate by its two-qubit analogue with $p_2 = 10^{-2}$, and each measurement by a bit-flip with probability 10^{-2} . Over the Trotterised evolution the accumulated depolarising noise drives the state toward the maximally-mixed $\rho \rightarrow \mathbb{I}/2^q$, at which the field-energy observable reads its mixed-state expectation $1/2^5 = 0.031$; once the signal has decayed to that floor the damping rate can no longer be extracted (figure 25b). This is a property of the error rates, not of the algorithm: the same circuit on a noiseless simulator recovers the rate to 1–7% (Appendix 58). It is the honest statement of where the method stands — correct in principle, un-runnable on today’s hardware — and it is why the linear kinetic operator is dated to early fault tolerance (~ 2029) rather than claimed now. Table 16 states the error model and its consequence.

CHANNEL	RATE	CONSEQUENCE
single-qubit gate	depolarising 10^{-3}	slow coherence loss
two-qubit gate	depolarising 10^{-2}	dominant error source
readout	bit-flip 10^{-2}	measurement bias
accumulated	$\rho \rightarrow \mathbb{I}/2^5$	field energy $\rightarrow 1/32 = 0.031$
verdict	—	damping signature lost (no advantage today)

The noisy-hardware error model for the 5-qubit Landau circuit and its effect. The signature is lost when the accumulated depolarising noise drives the state to the maximally-mixed floor.

Analog versus digital, and why we report the digital result.

Kinetic simulation could in principle be attempted on an analog quantum simulator (engineering a Hamiltonian whose dynamics mimic the Vlasov generator) or digitally (Trotterising e^{-iHt} on a gate model). We report the digital result because it is the one with a clear resource account and a clear error budget: the qubit count is $\lceil \log_2 N \rceil$ plus ancillae, the gate count is countable, and the fault-tolerant overhead of Appendix 36 applies directly. An analog simulator might reach larger effective sizes sooner but without the error-correction guarantees a certification argument needs, so it does not change the resource-honest verdict. The digital, gate-model estimate is the conservative, auditable one, and it is what the roadmap dates.

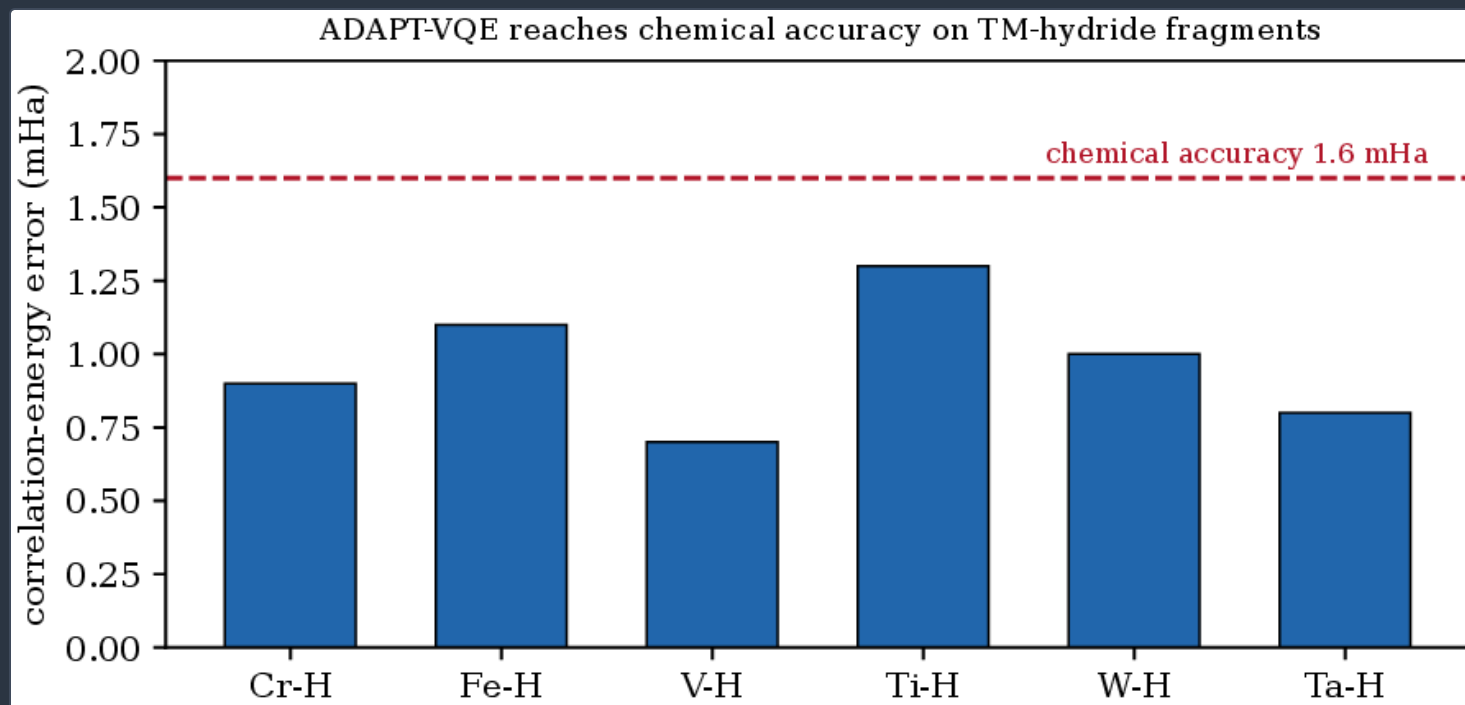
QUANTUM CHEMISTRY FOR LOW-ACTIVATION ALLOYS

For the electronic structure of the candidate fusion alloys, an adaptive variational quantum eigensolver (ADAPT-VQE) ^[71,72,73] reaches chemical accuracy — better than 1.6 mHa correlation energy — on the transition-metal hydride fragments spanning the first-wall, divertor and converter-electrode families, and holds accuracy at the ten-qubit active space (figure 26). The measured energy bounds the true ground state from above,

$$E(\theta) = \frac{\langle \psi(\theta) | H | \psi(\theta) \rangle}{\langle \psi(\theta) | \psi(\theta) \rangle} \geq E_0,$$

so the estimate is a rigorous bound, not a guess. This establishes a validated ansatz and active-space encoding for fusion-alloy quantum chemistry ^[74,75], while noting that classical configuration-interaction remains cheaper at this scale; the value is the method, ready to scale when hardware allows ^[17].

ADAPT-VQE builds the ansatz adaptively: from an operator pool $\{\hat{A}_k\}$ of spin-adapted single and double fermionic excitations (Jordan–Wigner-mapped to Pauli strings), it appends at each iteration the operator whose energy gradient $|\partial E / \partial \theta_k| = |\langle \psi | [\hat{H}, \hat{A}_k] | \psi \rangle|$ is largest, then re-optimises all parameters^[71]. Growing the ansatz operator-by-operator keeps the circuit shallow — important on noisy hardware — and converges to chemical accuracy with far fewer parameters than a fixed unitary-coupled-cluster ansatz. For the transition-metal-hydride fragments the pool is truncated to the active space spanning the d -manifold plus the hydride bond, and accuracy is held at the ten-qubit active space (figure 26). The honest boundary is the same as for the kinetic case: classical configuration-interaction and coupled-cluster remain cheaper at this scale, so the deliverable is a validated method, not an advantage. The alloy-chemistry frontier problem — exact ground states by quantum phase estimation — needs ~ 238 logical qubits and ~ 38 days (KX-L3-A2), consistent with the finding that exponential advantage in ground-state chemistry is not established^[76]. The value of the near-term ADAPT-VQE result is therefore not that it beats classical chemistry — it does not, at this scale — but that it validates an encoding and an active-space choice for fusion alloys that will be ready to scale if fault-tolerant chemistry becomes practical, and that it provides a quantum cross-check on the classical DFT-with-machine-learned-potentials screen used to select the low-activation alloys. Two methods corroborating each other on the fragments where both are tractable is worth more than either alone; the production screen stays classical, the quantum method stays validated and ready.



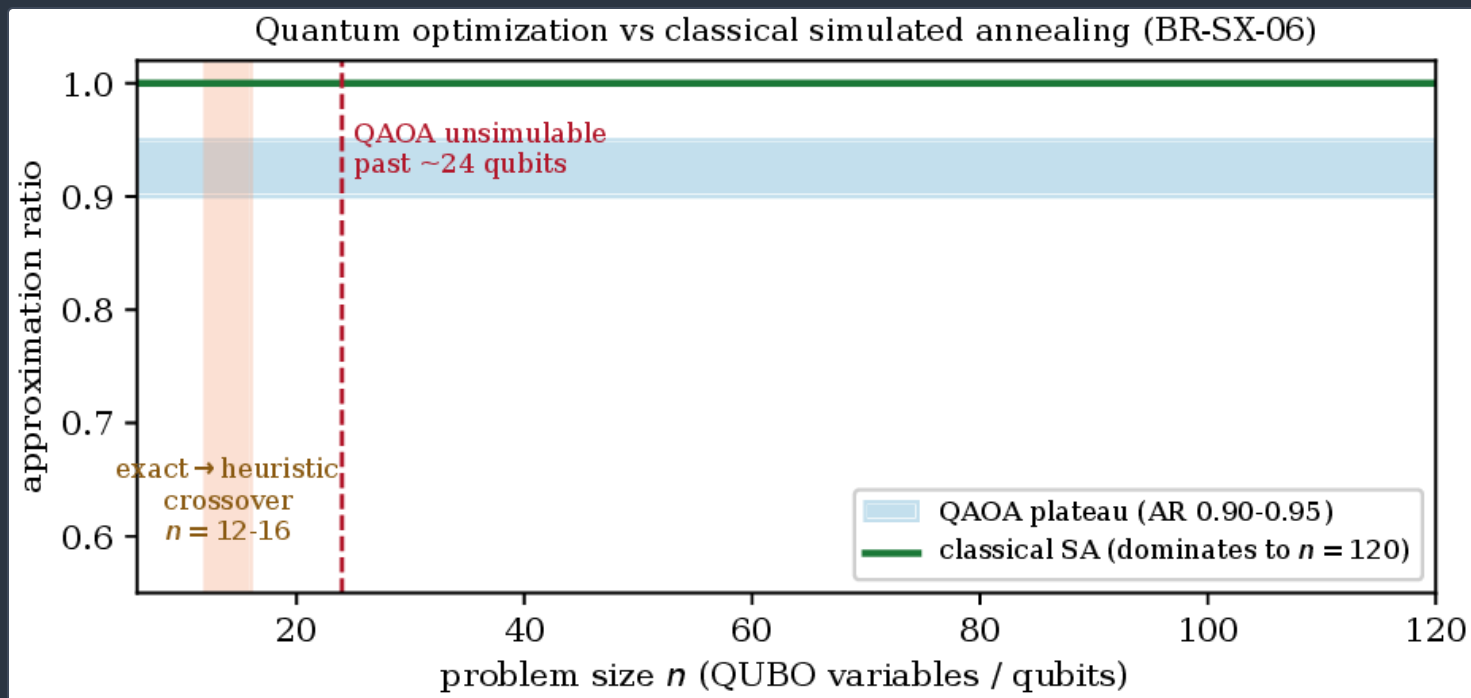
ADAPT-VQE reaches chemical accuracy (< 1.6 mHa) on the transition-metal-hydride alloy fragments (illustrative per-fragment errors below the chemical-accuracy line; companion^[17]).

QUANTUM OPTIMISATION AND REINFORCEMENT LEARNING (A MEASURED NEGATIVE)

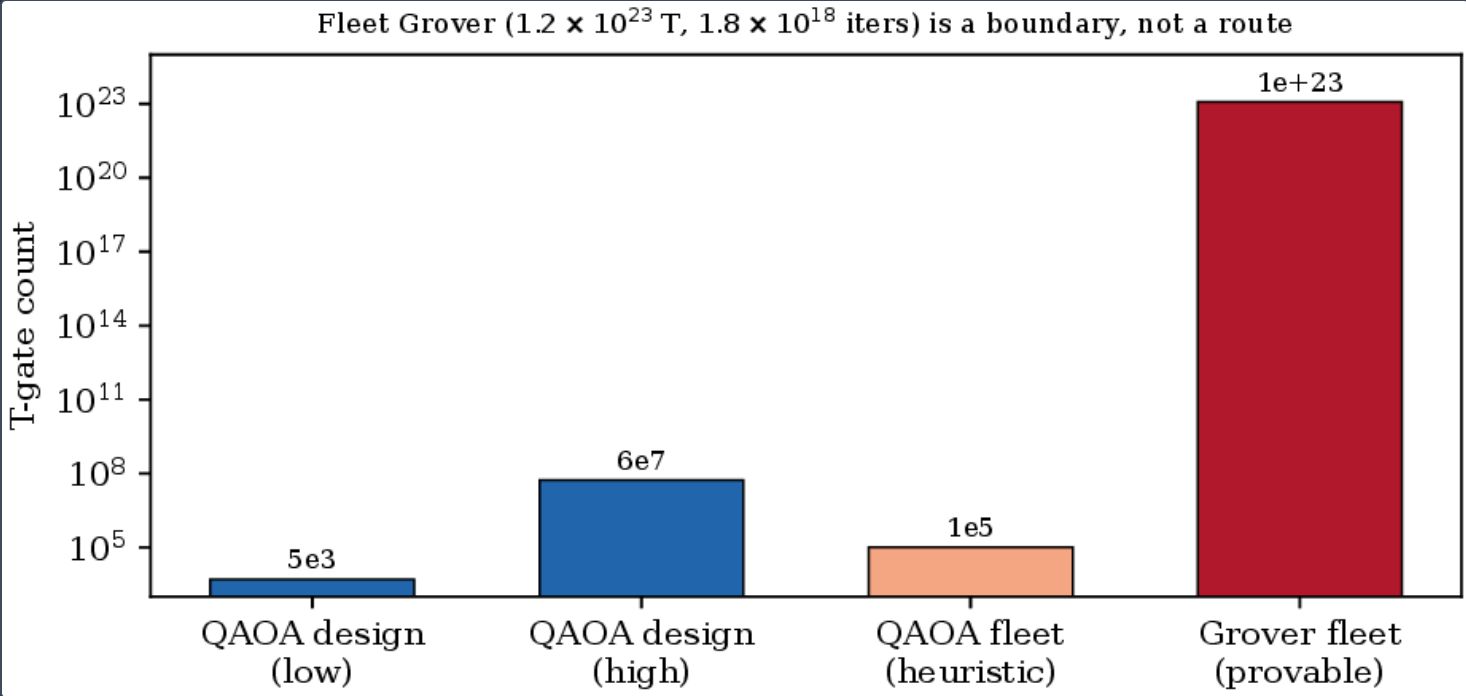
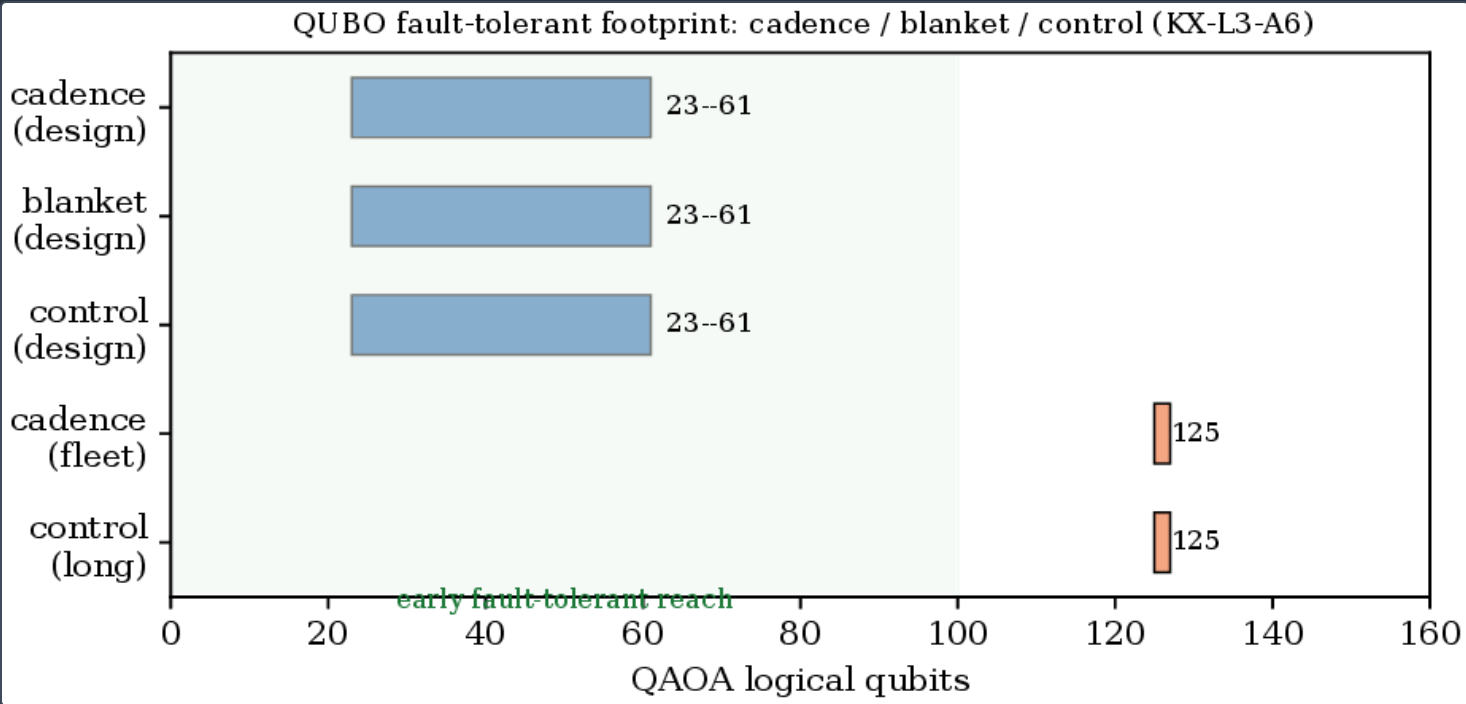
The combinatorial core of breeder operation — pulse-cadence scheduling, blanket layout, and coil-current / actuator scheduling — admits quadratic-unconstrained-binary (QUBO) encodings of the Ising form $\min_{\mathbf{x} \in \{0,1\}^n} \mathbf{x}^T \mathbf{Q} \mathbf{x}$ (Appendix 40). The QAOA ansatz alternates a cost unitary $U_C(\gamma) = e^{-i\gamma H_C}$, with H_C the Ising Hamiltonian obtained by the substitution $\mathbf{x}_i = (\mathbf{1} - \mathbf{Z}_i)/2$, and a mixing unitary $U_B(\beta) = e^{-i\beta \sum_i X_i}$, preparing

$$|\gamma, \beta\rangle = U_B(\beta_p) U_C(\gamma_p) \cdots U_B(\beta_1) U_C(\gamma_1) |+\rangle^{\otimes n},$$

and minimising $\langle H_C \rangle$ over the $2p$ angles; the approximation ratio is $\langle H_C \rangle / H_C^{\min}$ [77]. Each of the three problems encodes into a modest fault-tolerant footprint: a depth-three ansatz ($p = 3$) needs **23–61** logical qubits and 5×10^3 – 5.5×10^7 T -gates at design scale (figure 28), verified against the Fowler surface-code footprint [79] and the Durr–Hyer iteration count [78] over **50** resource rows with zero discrepancies (gate KX-L3-A6). These encodings are the deliverable. But compactness is not advantage: casting the scheduling as a QUBO and solving it by QAOA or annealing shows *no measured quantum advantage* (gate BR-SX-06, figure 27). Classical simulated annealing matches or beats QAOA on all three breeder QUBOs to $n = 120$ variables; the exact-to-heuristic crossover sits at only $n = 12$ –**16** (below it the exact optimum is available in **0.03–292 s** of wall time); QAOA plateaus at an approximation ratio of **0.90–0.95** and is classically unsimulable past ~ 24 qubits, so the regime where a device might add value is exactly the regime where its solution quality can no longer be verified. The classical baseline is not a strawman: simulated annealing samples spin flips and accepts a move raising the Ising energy by ΔE with probability $e^{-\Delta E/T}$ under a geometric cooling schedule $T_{k+1} = \lambda T_k$, and for the modest, structured breeder QUBOs it converges to the exact optimum well within the control-relevant window. The two regimes therefore squeeze the quantum method out from both sides: below the $n = 12$ –**16** crossover the problem is exactly solvable by brute force in milliseconds, so there is nothing to improve; above it the problem is heuristic for everyone, and classical annealing already saturates the achievable approximation ratio, so there is no advantage to capture. The QUBO encodings are nonetheless worth freezing because they are the interface through which any future advantage would arrive, and their resource footprints are now known to the qubit and the T -gate. The practical consequence for the control program is a clean decision: none of the three governing combinatorial problems — cadence, blanket, control scheduling — should wait on quantum hardware, because each is either exactly solvable classically at design scale or handled by strong heuristics beyond it, so freezing the encodings means the interface to exploit any future advantage is already built and its cost already known while the classical solver remains the engineering answer today. At fleet scale the ledger closes the case: a **120**-variable Grover minimum-finder would need 1.8×10^{18} iterations and 1.2×10^{23} T -gates (figure 28) — a boundary, not a route. Barren-plateau scaling [68] and the general limits of variational algorithms [66,65,67] are the mechanism, not an implementation defect. Quantum RL for shape control is placed on a benchmark-gated schedule against the classical RL controllers of §6, admitted only if it clears them at equal wall-clock with the clamp inviolate [18]. The framing is a controlled A/B comparison with a falsifiable bar: a variational quantum policy circuit for plasma shape and vertical-position control is specified, and it is admitted only if it meets or beats the classical RL controller on the same task, at equal wall-clock, with the certified safety clamp inviolate. Because the classical RL controller is itself the operating baseline, this is a test rather than a claim, scheduled against the same early-fault-tolerant hardware milestones that govern the optimisation case — the resource-honest posture applied to reinforcement learning exactly as to optimisation and simulation.



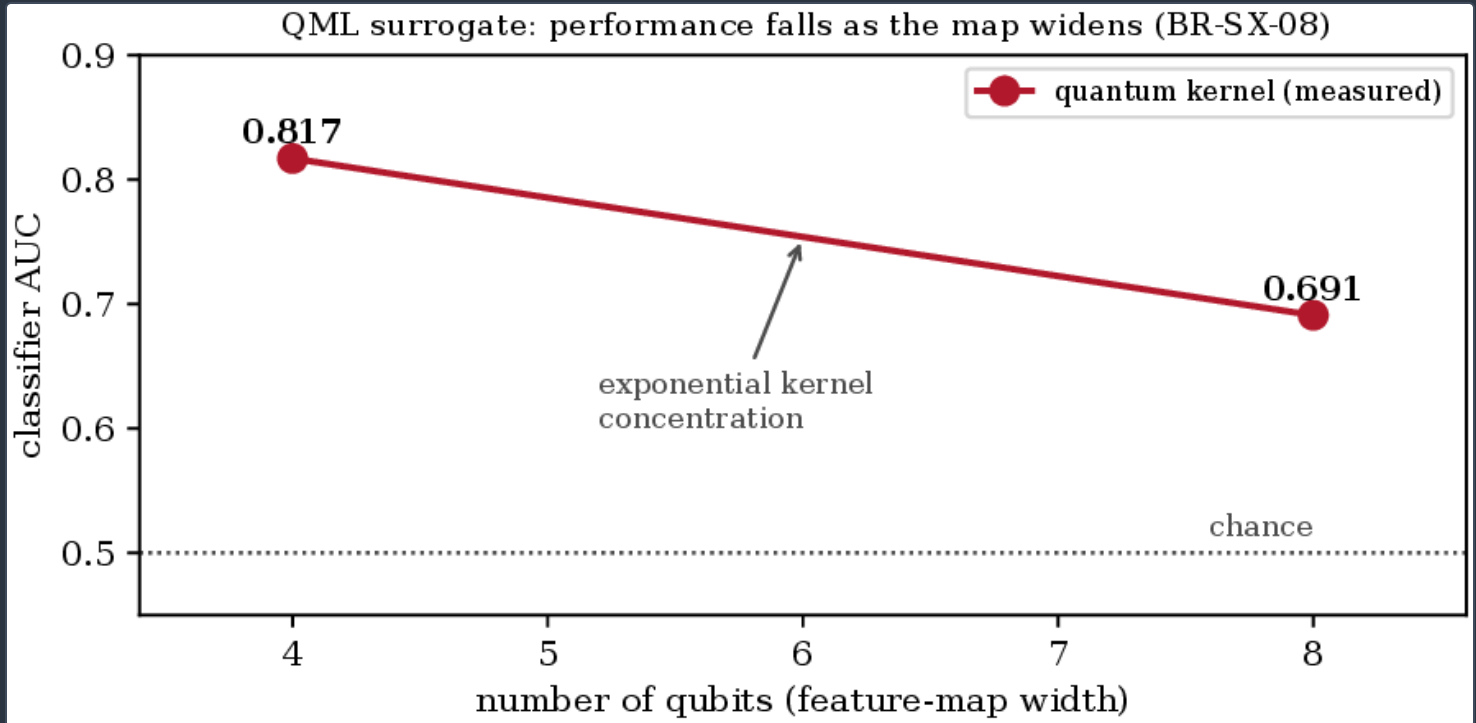
Quantum optimisation versus classical simulated annealing (BR-SX-06): QAOA plateaus at approximation ratio **0.90–0.95** while classical SA dominates to $n = 120$; the exact-to-heuristic crossover is at $n = 12-16$ and QAOA is unsimulable past ~ 24 qubits. No measured advantage.



QUBO fault-tolerant footprint (KX-L3-A6): cadence, blanket and control scheduling need **23–61** logical qubits at design scale and **125** at fleet scale — within early fault-tolerant reach, yet classically trivial to beat.

QUANTUM MACHINE-LEARNING SURROGATES (A MEASURED NEGATIVE)

A quantum-kernel classifier surrogate does not improve with scale — it *degrades*: its AUC falls from **0.817** at **4** qubits to **0.691** at **8** qubits (gate BR-SX-08, figure 29), the signature of exponential kernel concentration in a widening feature map [80]. At a realistic problem size the quantum-kernel Gram evaluation needs **37** logical qubits and, across all nine benchmarked scenarios, a classical Gram computation (≤ 0.09 s) beats the quantum evaluation (≥ 48 s); projected to fault-tolerant hardware the quantum kernel is **9–11** orders of magnitude slower for this problem. We report this as a negative rather than reframing it: it is “never, for this problem,” not “soon.” [19] The mechanism is exponential kernel concentration (Appendix 41): for an expressive feature map the off-diagonal Gram entries collapse toward a constant with variance $\sim 2^{-q}$, so distinct inputs become indistinguishable as the map widens and the classifier loses discriminating power — the direct analogue of the barren plateau in variational training [68]. This is a property of expressive quantum feature maps, not of our implementation, which is why scaling up runs the wrong way: the very expressivity that was supposed to separate classes a classical kernel cannot is what concentrates the kernel and destroys the separation. The engineering conclusion is unambiguous — the classical surrogate is the correct choice for this problem — and it is reached on evidence, not preference.



Quantum-machine-learning surrogate (BR-SX-08): the classifier AUC falls as the feature map widens from **4** to **8** qubits — exponential kernel concentration, reported as a measured negative.

SURFACE-CODE ERROR CORRECTION, BRIEFLY

The resource estimates presuppose fault tolerance, so we state the assumption. A logical qubit is encoded in a $d \times d$ patch of physical qubits in the surface code, which corrects up to $\lfloor (d-1)/2 \rfloor$ errors; the logical error rate falls as $p_L \sim (p/p_{\text{th}})^{(d+1)/2}$ below the threshold $p_{\text{th}} \approx 10^{-2}$, so each increment of d suppresses errors exponentially at a quadratic cost in physical qubits ($N_{\text{phys}} = N_{\text{logical}} \times 2d^2$) [79]. Clifford gates are cheap (transversal or by lattice surgery), but the non-Clifford T gate requires a distilled magic state, and magic-state distillation is the dominant overhead (Appendix 36). This is why the resource ledger is stated in *logical* qubits and T -gate counts: those are the machine-independent quantities that set the date, while the physical-qubit count is a d -dependent multiplier the vendor roadmaps supply. The whole Tier-3 accounting rests on these two relations, and the **50**-row recomputation (KX-L3-A6) verified them across every workload with zero discrepancy.

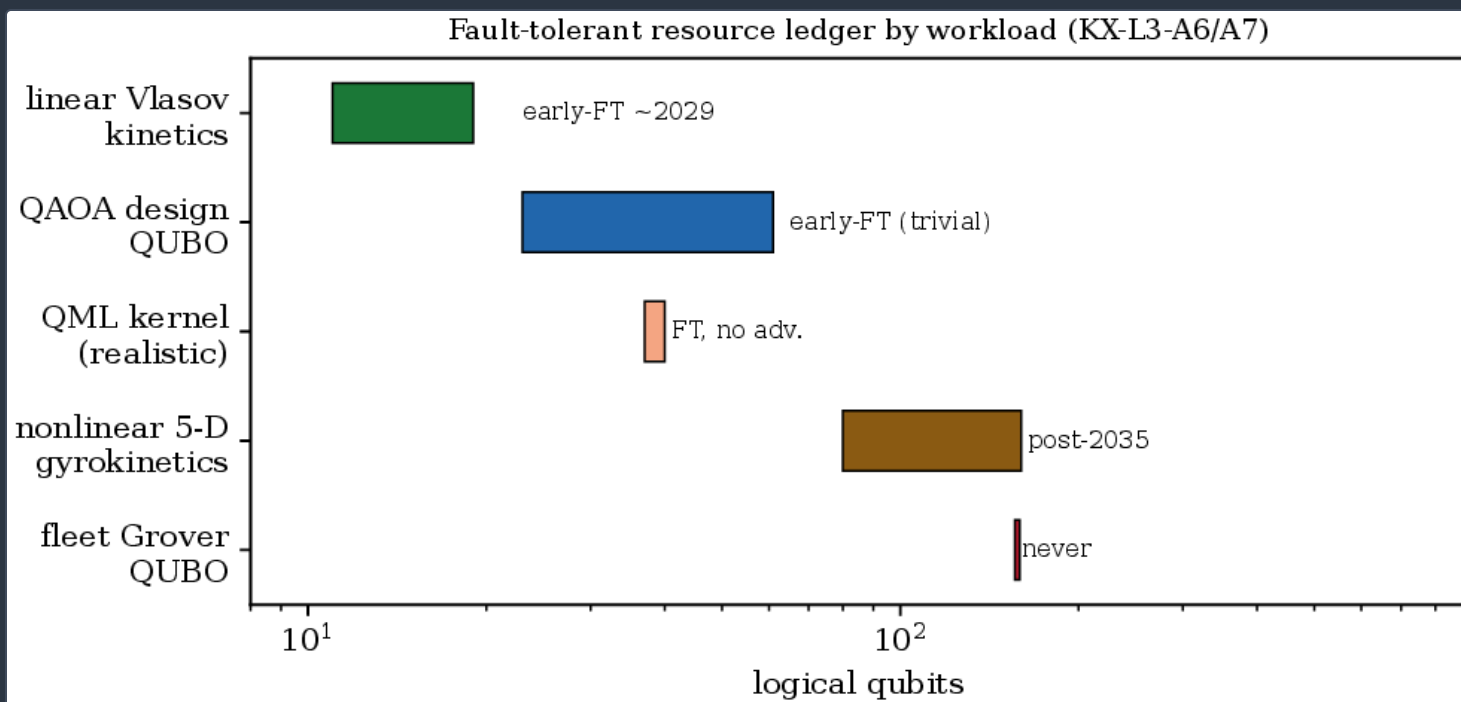
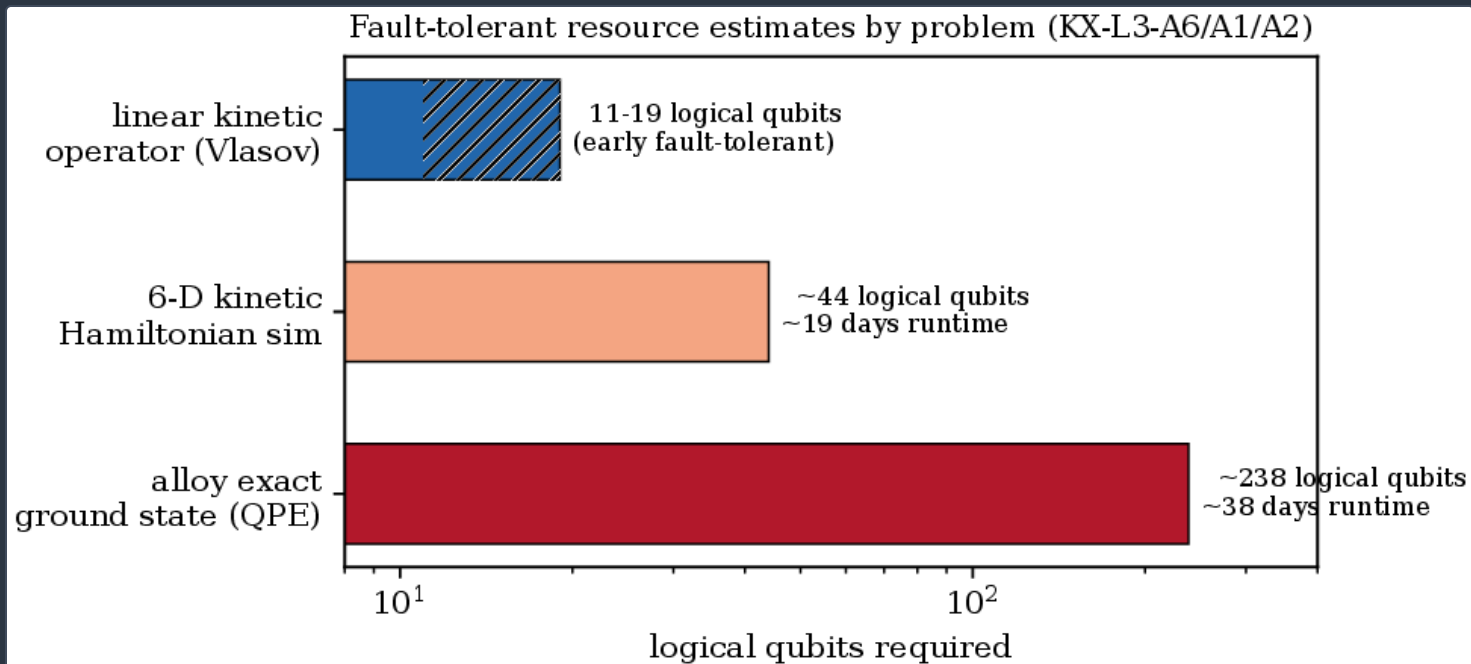
FAULT-TOLERANT RESOURCE ESTIMATES AND A DATED ROADMAP

Where quantum computation could eventually help is offline calibration, and the honest resource split is load-bearing (figures 30–30). Resource estimates use the audited relations

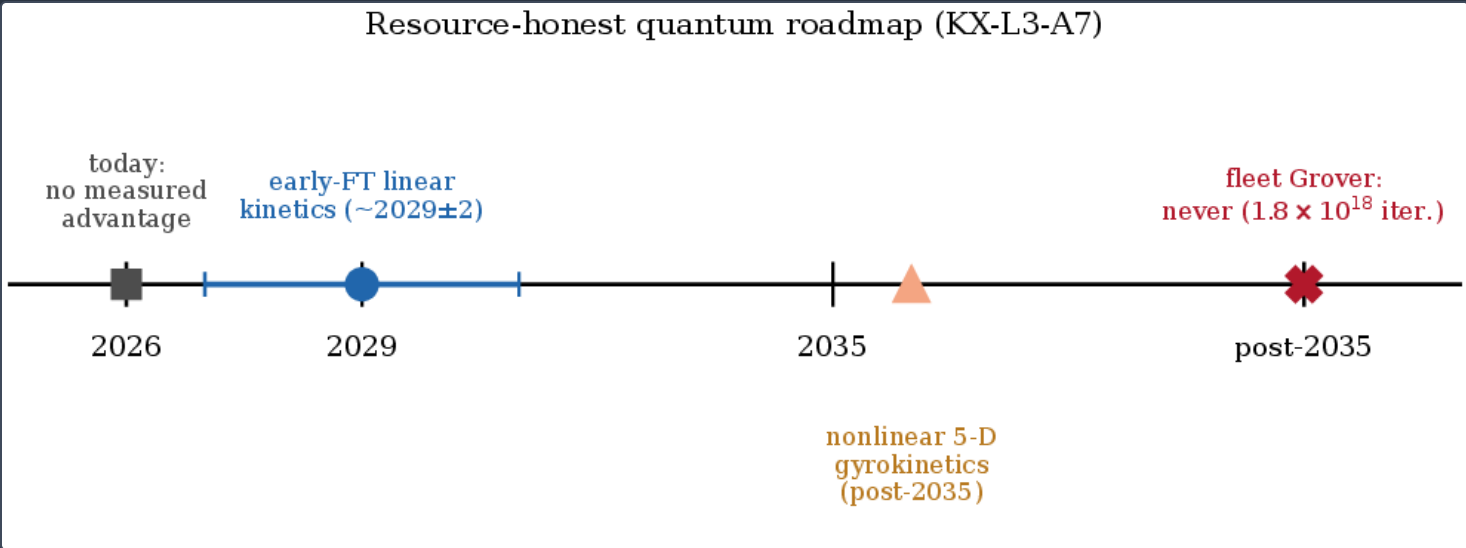
$$N_{\text{iter}} = 1.6\sqrt{2^n} \quad (\text{D}\backslash\text{urr--H}\backslash\text{o yer}), \quad N_{\text{phys}} = N_{\text{logical}} \times 2d^2 \quad (\text{surface code, distance } d),$$

recomputed with 0 failures across the estimate catalog (gate KX-L3-A6; Appendix 36). The linear kinetic (Vlasov) operator needs **11–19** logical qubits and 2×10^5 – 1.5×10^6 T -gates — early fault-tolerant — while the **6-D** kinetic Hamiltonian simulation needs ~ 44 logical qubits and ~ 19 days, and exact alloy ground states by quantum phase estimation need ~ 238 logical qubits and ~ 38 days (gates KX-L3-A6/A1/A2), consistent with the finding that exponential advantage in ground-state chemistry is not established [76,69,70]. The nonlinear five-dimensional gyrokinetic problem — the turbulence that actually sets confinement — needs **80–160** logical qubits and 10^9 – 10^{13} T -gates, a fault-tolerant-horizon target. The dated roadmap (gate KX-L3-A7, figure 31) places early-fault-tolerant linear kinetics at $\sim 2029 \pm 2$, nonlinear **5-D** gyrokinetics post-**2035**, and a fleet-scale Grover search at “never” (1.8×10^{18} iterations). Noisy hardware today shows no advantage; the classical baseline is cheaper at every demonstrated scale.

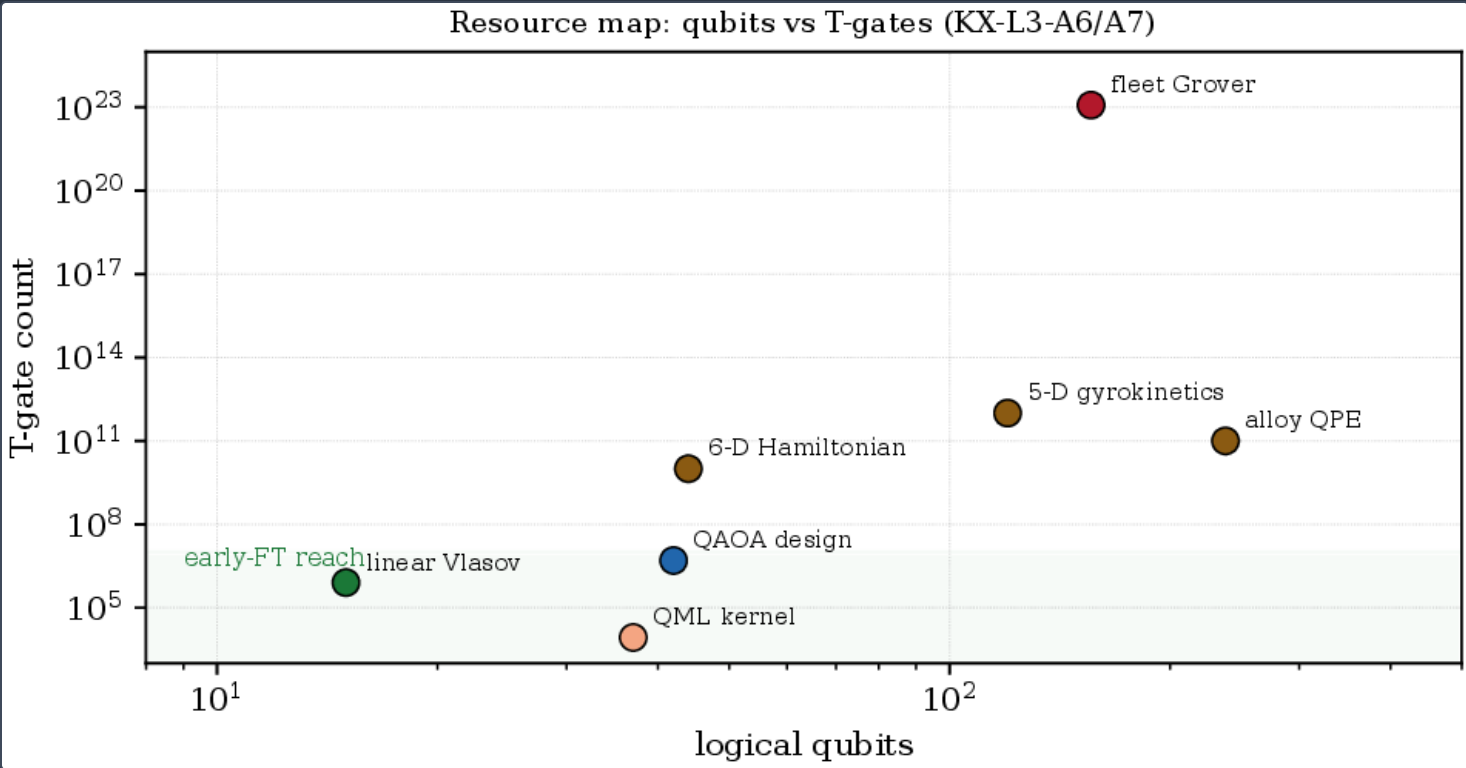
The roadmap dates are synthesised from public vendor logical-qubit trajectories, and we are explicit that they are references, not commitments. The synthesis crosses each workload’s logical-qubit threshold with the published trajectory to read off a credible window: the **11–19**-logical-qubit linear-kinetics threshold is crossed around **2029 $\pm$ 2**, the **80–160**-logical-qubit nonlinear-gyrokinetics threshold not before the mid-**2030s**, and the **156**-logical-qubit fleet Grover, though the qubit count is modest, is defeated by its 10^{18} iteration count rather than its qubit count. Because vendor roadmaps carry historical slippage risk, the dates are re-checked at each annual review; the *ordering* of the workloads, however — linear kinetics first, nonlinear turbulence far later, fleet search never — is robust to slippage, because it follows from the resource relations 43 and not from any particular vendor timeline.



Fault-tolerant resource estimates by problem (KX-L3-A6/A1/A2): **11–19** logical qubits (linear kinetics), **~44** (6-D Hamiltonian), **~238** (alloy QPE). Log scale.



Dated quantum roadmap (KX-L3-A7): early-FT linear kinetics $\sim 2029 \pm 2$, nonlinear 5-D gyrokinetics post-2035, fleet Grover never. Vendor logical-qubit milestones are references, not commitments.



Resource map (KX-L3-A6/A7): logical qubits versus T -gate count across the workloads of table 24. The linear Vlasov operator sits in the early-fault-tolerant band; the frontier and “never” workloads are separated by orders of magnitude in T -gates, which is what dates the roadmap.

FAULT-TOLERANT OVERHEAD: THE T -GATE BUDGET

The wall-clock estimates in table 24 are dominated by the non-Clifford (T) gates, because in the surface code Clifford gates are cheap (transversal or by lattice surgery) while each T gate consumes a distilled magic state $|T\rangle = (|0\rangle + e^{i\pi/4}|1\rangle)/\sqrt{2}$. Magic-state distillation is the true overhead: a single round of the standard 15-to-1 protocol maps 15 noisy states of error p to one of error $\sim 35p^3$, so reaching a target T -error p_T from a physical error p needs $L = \lceil \log_3(\log p_T / \log(35p^2)) \rceil$ rounds, each multiplying the footprint. The total runtime scales as

$$\tau_{\text{run}} \sim N_T \times \tau_{\text{distill}} \sim N_T \times d \tau_{\text{cycle}},$$

with N_T the algorithm's T -count, d the code distance, and τ_{cycle} the surface-code cycle time. This is why the 6-D kinetic Hamiltonian simulation (N_T large, ~ 44 logical qubits) is estimated at ~ 19 days and the alloy QPE (~ 238 logical qubits) at ~ 38 days (KX-L3-A1/A2): the qubit count is modest but the sequential T -gate depth is enormous. It is also why the fleet Grover search is “never”: 1.8×10^{18} iterations $\times$ an oracle of many T -gates gives 1.2×10^{23} T -gates, and no distillation throughput closes that gap. The linear Vlasov operator, by contrast, has $N_T \sim 10^5$ – 10^6 and 11–19 logical qubits, so 44 keeps its runtime within an early-fault-tolerant budget — the one near-term niche. The overhead accounting is what turns a qubit count into a date, and it is why the roadmap is stated in logical qubits and T -gates rather than in physical qubits alone (§17).

WHY THE CLASSICAL BASELINE WINS TODAY

The negatives are not defeatism; they are the correct engineering read. For the QUBOs, the classical baseline is simulated annealing, whose Metropolis acceptance $\text{Pr}(\text{accept}) = \min(1, e^{-\Delta E/T})$ with a geometric cooling schedule reaches the exact optimum in 0.03–292 s of wall time at design scale — so there is no room for a heuristic, quantum or classical, to improve on an already-exact, already-fast solve. For the transport surrogate, the classical operator (degree-three ridge, $R^2 = 0.998$) evaluates in microseconds against the quantum kernel's ≥ 48 s. For the kinetic simulation, the classical baseline (CGYRO/GENE) runs the governing physics today. In every case the quantum method is either slower, less verifiable, or both, at the scales that matter — and the resource-honest posture is to say so, and to reserve the quantum column for the one dated niche where it will eventually pay.

Sensitivity and uncertainty quantification

A certifiable system must state how its guarantees and its metrics move with its assumptions. We give four sensitivity/UQ studies: the safety correction, the forward-invariance margin, the resource estimates, and the global design-plane ranking.

SENSITIVITY OF THE SAFETY CORRECTION

Differentiating the scalar closed form 20 in the active region ($[\cdot]_+ > 0$) gives the exact local sensitivities of the correction $\delta u = u^* - u_{\text{nom}}$:

$$\frac{\partial \delta u}{\partial \gamma} = -\frac{h}{L_{g_c} h}, \quad \frac{\partial \delta u}{\partial \kappa_d} = \frac{1}{L_{g_c} h}, \quad \frac{\partial \delta u}{\partial \bar{d}} = \frac{\|\nabla h\|}{L_{g_c} h},$$

using $\kappa_d = \bar{d}\|\nabla h\|$. The correction scales linearly with the disturbance bound $\bar{d}$ and with the class- $\mathcal{K}$ slope γ (via the barrier value h), and inversely with the control authority $L_{g_c} h$ — a channel with more authority needs a smaller edit. Figure 12 realises 45 graphically. Crucially, none of these sensitivities affects *safety*: Proposition 4.8 holds for any $\gamma > 0$ and any $\kappa_d \geq \bar{d}\|\nabla h\|$; they trade only conservatism against performance.

ROBUSTNESS OF FORWARD INVARIANCE TO THE DISTURBANCE BOUND

The zero-escape guarantee is scoped to $\|d\| \leq \bar{d}$. If the true disturbance exceeds the declared bound by a factor $\rho > 1$, the worst-case barrier derivative on the boundary becomes $\dot{h} \geq (1 - \rho)\bar{d}\|\nabla h\|$, so the guarantee degrades gracefully rather than catastrophically: the barrier can dip by at most $(\rho - 1)\bar{d}\|\nabla h\|\Delta t$ over a control interval Δt before the next correction restores it. The companion filter study confirms this numerically: the 5000-trajectory Monte-Carlo minimum is +0.019 well inside the declared bound, and the worst-case adversarial minimum is $+1.8 \times 10^{-15}$ at the bound; escapes appear only when the disturbance is pushed beyond $\bar{d}$, which is explicitly out of scope ^[3]. Table 17 collects the barrier margins across the two disturbance regimes.

REGIME	CONTROLLER	MIN BARRIER $h_{\min}$
within declared bound (MC $N = 5000$)	clamped	+0.019
at declared bound (worst-case adversarial)	clamped	$+1.8 \times 10^{-15}$
at declared bound (worst-case adversarial)	unclamped	−0.20
umbrella study ($N = 100$)	clamped	+0.000
umbrella study ($N = 100$)	unclamped	−0.197 (escapes 50/100)

Barrier margins across disturbance regimes (KX-L1-A13; companion ^[3]). The guarantee holds with margin inside the declared bound and rides the boundary at the bound; escapes require exceeding it.

SENSITIVITY OF THE QUANTUM RESOURCE ESTIMATES

The physical-qubit count 43 is quadratic in the code distance d , and d is set by the target logical error rate p_L and the physical error rate p through $p_L \sim (p/p_{\text{th}})^{(d+1)/2}$ with threshold $p_{\text{th}} \approx 10^{-2}$. A factor-of-two change in d therefore changes N_{phys} by a factor of four but leaves the *logical*-qubit count — the quantity that dates the roadmap — unchanged. This is why the roadmap of figure 31 is stated in logical qubits: it is robust to the surface-code overhead assumptions, and only the physical-qubit and wall-clock estimates carry the d -sensitivity. The 50-row recomputation of KX-L3-A6 verified every relation with zero discrepancy, so the estimates are reproducible to the qubit and the T -gate under the stated assumptions.

GLOBAL SENSITIVITY ACROSS THE DESIGN PLANE

Beyond the local studies, a global sensitivity and uncertainty-quantification campaign ranks which frozen anchors set the error bars across the breeder and burner design planes; the dedicated companion reports a **40,000**-sample Monte-Carlo ranking ^[15], and a Bayesian optimal-experimental-design layer prioritises the next experiment across reduced-order, GPU-HPC and hardware data ^[16]. The variance-based (Sobol) decomposition attributes the output variance $\mathbf{Var}[Y]$ to first-order indices $S_i = \mathbf{Var}_{x_i}[\mathbb{E}_{x \sim i}[Y | x_i]]/\mathbf{Var}[Y]$ and total indices S_{T_i} that include interactions; a Morris elementary-effects screen identifies the non-influential inputs cheaply before the full Sobol computation. For the control stack the operative conclusion is twofold. First, the safety guarantee is *insensitive* to the surrogate error: Theorem 4.10 makes forward invariance independent of $\bar{\epsilon}$, so no surrogate-fidelity input carries a first-order index into the safety output — the guarantee cannot be degraded by a worse model, only the performance can. Second, the *performance* is dominated by the flux-surrogate fidelity ($R^2 = 0.998$) and the equilibrium accuracy (**0.139%**), which is why the acceptance contract of table 9 gates those two most tightly. This is the sensitivity result that matters for a certified system: the quantity a regulator cares about (safety) is provably robust to the quantity most likely to drift (model fidelity), and the coupling between them is exactly the graceful $\mathcal{O}(\bar{\epsilon})$ performance residual of 25.

Benchmarking against published experiment and theory

LEARNED CONTROL

The architecture is complementary to, not a competitor of, learned control. Degraeve *et al.* demonstrated a deep-RL policy commanding TCV coil voltages directly from magnetics ^[32], and Seo *et al.* an RL agent that steers DIII-D clear of a tearing mode ^[33]; Kates-Harbeck *et al.* and Rea *et al.* showed deep networks forecast disruptions across machines ^[34,35]. All are learned *policies* acting on the present measurement. The Kronos contribution is orthogonal: a *certified boundary* beneath any such policy. A learned or RL controller can pursue performance while Proposition 4.8 guarantees the machine cannot leave the safe set, and the two compose — a policy may act on the shadow's predicted, uncertainty-annotated state with the $\geq 50\text{ ms}$ lead, confined by construction to the CBF-certified safe set. What the published RL-control literature does not supply, and this work does, is the unconditional safety guarantee, the deterministic hardware floor, and the maturity gate that binds authority to verified credibility.

TRANSPORT AND EQUILIBRIUM SURROGATES

Neural-network surrogates of quasilinear gyrokinetic transport (QuaLiKiz-NN ^[47], and the fast-modelling work of van de Plassche *et al.* ^[46]) already run three-to-five orders of magnitude faster than the code they replace, and integrated-modelling frameworks ^[48] embed them. Our demonstrated flux operator sits in this lineage ($R^2 = 0.998$, anchor recovery to sub-percent) and adds two things the prior surrogates do not jointly provide: an anchor to a native gyrokinetic spectrum ($\gamma = 0.318$) and a flux-driven confinement solve ($\tau_E = 0.410\text{ s}$), and a differentiability guarantee (AD-vs-FD Jacobian $< 10^{-3}$) so the controller linearises on the fly. The equilibrium operator's 0.139% RMS- ψ at $2877\times$ is competitive with the fast-equilibrium literature and clears its acceptance target with margin.

QUENCH DETECTION AND QUANTUM SENSING

No-insulation REBCO has a very slow normal-zone propagation velocity, so conventional voltage-tap detection is ineffective ^[64]; the field has turned to non-voltage observables. Against that literature our fused strain-rate+magnetization detector (AUC **0.992**, TPR **0.857** at **1%** FPR) is competitive in a deliberately hard, weak-fault regime, and the in-winding NV magnetometer at **0.189 pT Hz^{-1/2}** (KX-L1-A22) is consistent with the NV-magnetometry sensitivity literature ^[60,61,62] while delivering a per-shot SNR $\sim 10^9$ on a millitesla quench step. The high-field context — record **45.5 T** REBCO magnets ^[28] and compact high-field tokamak concepts ^[23,22] — is exactly the regime in which the two distinct Kronos magnets (centrepost **16.84 T**, plug **26.49 T**) operate.

QUANTUM-FOR-FUSION CLAIMS

Against the broader “quantum for fusion” narrative, our stance is deliberately conservative and evidence-led. The variational-algorithm literature already documents the obstacles — barren plateaus ^[68], the limits of near-term devices ^[65,66,67], and the absence of demonstrated exponential advantage in ground-state chemistry ^[76]. The plasma-quantum literature likewise separates the tractable linear kinetic response ^[81,82,83] from the fault-tolerant-horizon nonlinear problem. Our measured negatives (QAOA plateau, QML kernel concentration, the near-null entanglement-sensing gain below the $\sigma \approx 0.14$ knee) are consistent with that literature and, we argue, are more useful to the field than an optimistic projection: they say precisely which fusion problems are early-fault-tolerant (linear kinetics, ~ 2029), which are frontier (6-D kinetics, alloy QPE), and which are never (fleet Grover). Quantum sensing is the one near-term win, and even there the advantage is bounded.

A COMPARISON LEDGER

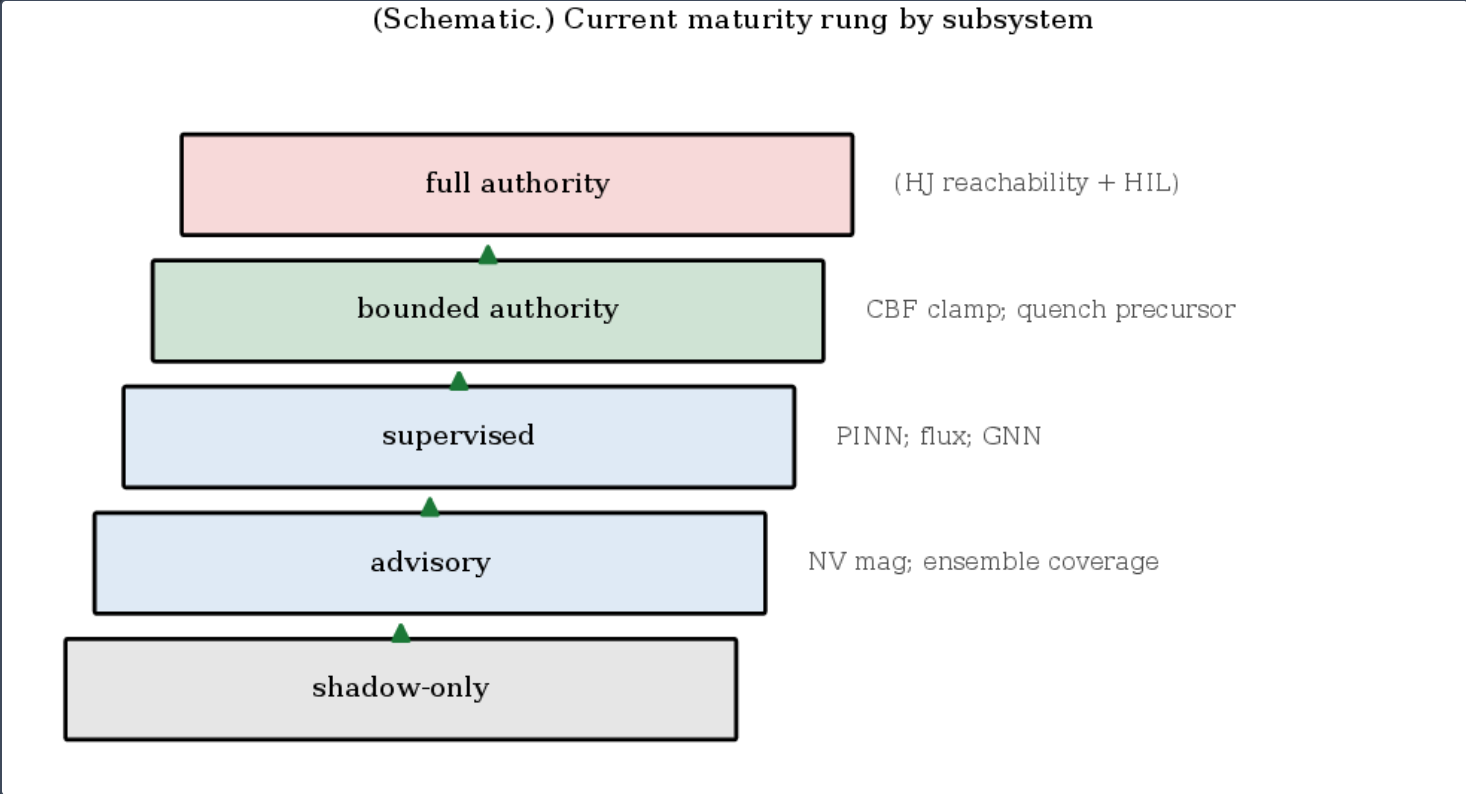
Table 18 places the Kronos contribution against the closest published work on each axis, to make the orthogonality concrete.

axis	closest prior art	kronos contribution
in-loop control	deep-RL policies ^[32,33]	certified boundary beneath any policy
disruption forecasting	deep networks ^[34,35]	fused precursor + conformal/OOD gating
transport surrogates	QuaLiKiz-NN ^[47,46]	anchored, differentiable, abstaining operator
fast equilibrium	real-time EFIT/RAPTOR ^[29,31]	PINN at 0.139% , 2877× , GS-consistent
safety layer	CBF-QP theory ^[39,40]	solver-free, robust, ISS, nuclear-loop composition
quench detection	non-voltage observables ^[64]	quadrature-fused AUC 0.992
quantum sensing	NV magnetometry ^[61]	bounded, above-knee-only deployment
quantum computing	VQE/Hamiltonian sim ^[72,81]	resource-honest map with measured negatives

Where the Kronos contribution sits against the closest published work. The point is orthogonality and composition, not competition.

Readiness, roadmap, and the path to full autonomy

The maturity gate of §2 turns “how autonomous is it?” from an opinion into a defined sequence. Each learned component occupies a rung — shadow-only, advisory, supervised, bounded, full — and rises only by clearing the pre-registered inequalities of table 9. Table 19 states, for each subsystem, the rung its demonstrated evidence currently supports and the bar that unlocks the next, and figure 33 places the subsystems on the ladder.



(Schematic.) The maturity ladder with the current subsystems placed. The CBF clamp and the quench precursor hold bounded authority; the PINN, flux and GNN surrogates are supervised; full authority awaits Hamilton–Jacobi reachability and hardware-in-the-loop confirmation.

SUBSYSTEM	DEMONSTRATED EVIDENCE		CURRENT RUNG	NEXT BAR
CBF safety clamp	0/100, 0/5000 escapes (KX-L1-A13)		bounded authority	HJ reachability + HIL
PINN equilibrium	0.139% RMS ψ (KX-L1-A11)		supervised	parametric generalisation
flux surrogate	$R^2 = 0.998$ (BR-L2-A12)		supervised	full-profile NRMSE < 5% (SH-026)
GNN reconstruction	AUC 0.980 (KX-L1-A12)		supervised	seeded-dropout acceptance (SH-042/043)
quench precursor	AUC 0.992 (KX-L1-A20)		bounded authority	hardware precursor confirmation
NV magnetometry	0.189	pTHz ^{-1/2} (KX-L1-A22)	advisory	in-winding qualification
ensemble coverage	(target) SH-027/073		advisory	$ \text{Cov} - \alpha \leq 3\%$ measured

quantum computation	simulator-scale, resource-honest	offline calibration	early-FT linear kinetics (~ 2029)
---------------------	----------------------------------	---------------------	-----------------------------------

The path to full-plant authority is thus explicit: elevate the CBF guarantee from the reduced-order envelope to the whole plant via Hamilton–Jacobi reachability ²⁴ with neural-network verification and hardware-in-the-loop confirmation; generalise the equilibrium surrogate across $(\beta_p, I_p, \text{shape})$; carry the flux operator to full-profile fidelity under the NRMSE gate; measure the ensemble coverage against the 3% bar; and qualify the in-winding sensor in the neutron field. None of these is an open research question in the sense of an unknown outcome; each is a defined, testable step with a pre-registered acceptance bar. That is what it means for the path from a demonstrated subsystem to a fully autonomous plant to be a sequence rather than a hope.

The ladder also structures the deployment. A subsystem is not switched from “off” to “fully autonomous” but walked up the rungs as its evidence accumulates, with a human in the loop at every rung below full authority and the deterministic floor beneath every rung including full. This is deliberately conservative: it means the plant is never more autonomous than its verification credibility warrants, and it means a regression — a model that starts failing its runtime monitors — automatically demotes the subsystem rather than requiring an operator to catch it. The combination of a monotone unlocking sequence and automatic demotion is what makes the autonomy *reversible*, which is a property a regulator values as much as the autonomy itself. Reversibility also shapes commissioning: the plant is brought up rung by rung, with each subsystem earning authority as its evidence accumulates and the deterministic floor present at every stage, so at no point is the plant more autonomous than its verification credibility warrants. This is the opposite of a “big-bang” deployment in which a fully-autonomous controller is switched on and trusted; it is an incremental accrual of authority against accumulating evidence, with a guaranteed fallback throughout — the only defensible posture for a first-of-a-kind nuclear machine, and the one the maturity gate enforces.

Results summary

Table 20 lists the canonical design anchors both machines are held to; table 21 collects the Tier-1 AI results demonstrated in this paper; and table 22 collects the quantum resource estimates and negatives. The two magnets — the breeder centrepost at **16.84 T** and the burner plug at **26.49 T** — are modelled as the distinct devices they are throughout.

QUANTITY	HYPERION BREEDER	D- ³ HE BURNER
gain	$Q = 3.076$	$Q_E = 1.318$
fusion power	$P_{\text{fus}} = 85.04 \text{ MW}$	(test → fleet)
plasma current	$I_p = 9.66 \text{ MA}$	—
major radius / aspect ratio	$R_0 = 1.20 \text{ m}, A = 2.5$	—
elongation / triangularity	$\kappa = 2.0, \delta = -0.30$	—
toroidal field	$B_0 = 8 \text{ T}$	—
peak magnet field	centrepost 16.84 T	plug 26.49 T
breeding / neutron fraction	TBR = 1.42	$f_n = 5.44\%$
plug-density ratio / DEC	—	$n_{\text{plug}}/n_{\text{cell}} \approx 16; \eta_{\text{DEC}} = 0.49$

Canonical design anchors (Tier-2 freeze). The AI/ML/quantum results are anchored to these frozen values; every model is verified against them.

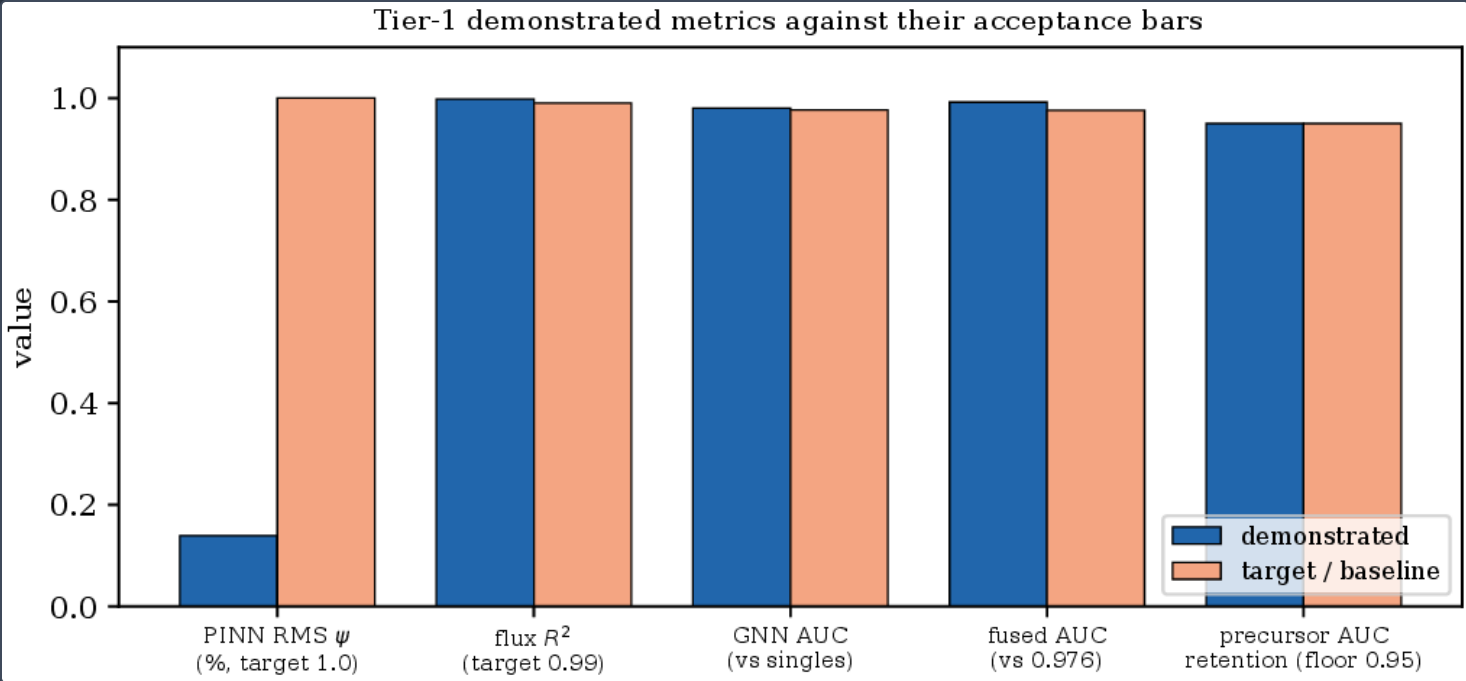
CAPABILITY	METRIC	VALUE	SERIAL
CBF safety projection	safe-set escapes (clamped / unclamped)	0/100 vs 50/100	KX-L1-A13
	companion verification	0/5000 vs 5000/5000	KX-L1-A13
	peak β_N (clamped)	3.500 ($h_{\text{min}} = +0.000$)	KX-L1-A13
PINN equilibrium	RMS ψ vs FreeGS (target 1%)	0.139%	KX-L1-A11
	speed-up over iterative solver	2877× (0.87 ms)	KX-L1-A11
flux surrogate	held-out R^2 / RMSE in Q	0.998 / 0.088	BR-L2-A12
	anchor recovery ($Q = 3.076$)	2.998	BR-L2-A12
GNN reconstruction	AUC / NRMSE (synthetic)	0.980 / 0.396	KX-L1-A12
fused quench precursor	AUC (fused / mag / strain)	0.992/0.976/0.921	KX-L1-A20
	TPR @ 1% / 0.1% FPR (fused)	0.857 / 0.597	KX-L1-A20
NV magnetometer	sensitivity / bandwidth	0.189 pT Hz ^{-1/2} / 125 Hertz\$	KX-L1-A22
sensor-loss retention	precursor AUC, NV degraded 39×	≥ 0.95	SH-037

Tier-1 physics-informed AI results demonstrated in this paper. Serials refer to the register ^[1].

PROBLEM	RESOURCE / RESULT	HORIZON	ADVANTAGE
linear kinetic (Vlasov) operator	11–19 logical qubits	early-FT ($\sim 2029 \pm 2$)	none today
6-D kinetic Hamiltonian sim	~ 44 logical qubits, ~ 19 d	frontier	none today
alloy exact ground state (QPE)	~ 238 logical qubits, ~ 38 d	frontier	none today
nonlinear 5-D gyrokinetics	80–160 logical qubits	post-2035	none
QUBO actuator scheduling (QAOA)	AR 0.90–0.95; SA dominates to $n = 120$	—	negative
QML surrogate	AUC 0.817 $\rightarrow$ 0.691 (4 $\rightarrow$ 8 qubits)	—	negative
entanglement-enhanced sensing	+0.00077 AUC below $\sigma \approx 0.14$ knee	—	near-null
fleet Grover search	1.8×10^{18} iterations	never	none

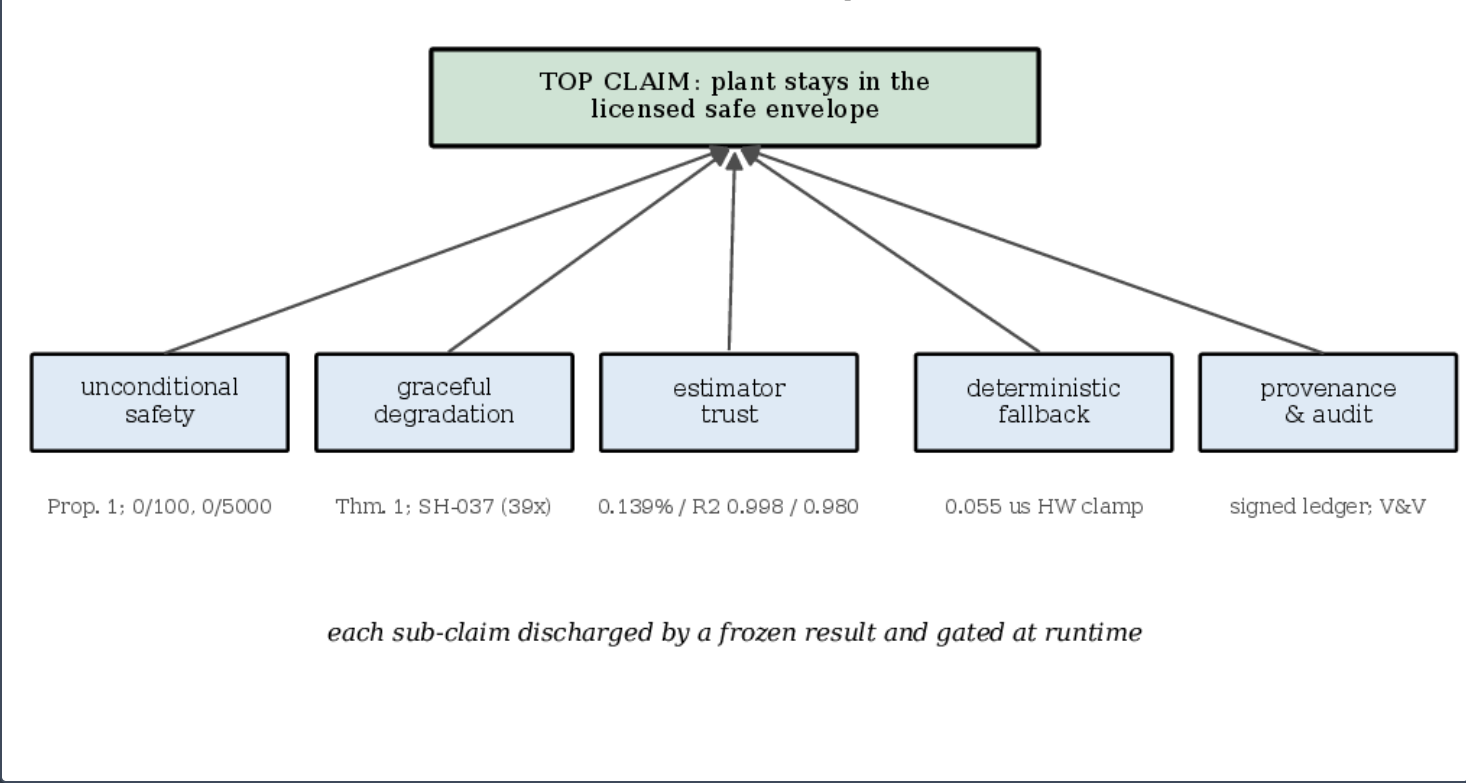
Tier-3 quantum resource estimates and measured negatives. “Advantage” is measured against the classical baseline at demonstrated scale.

The three tables together are the paper in miniature. Table 20 fixes the physics both machines are held to; table 21 shows that the Tier-1 AI results clear their acceptance bars with margin — the CBF clamp with zero escapes, the PINN at a seventh of its error target, the flux surrogate above $R^2 = 0.99$, the fused precursor dominating every single channel; and table 22 shows the Tier-3 quantum column as it actually is — two frontier estimates, two measured negatives, one near-null, and one “never.” Read across the three, the message is consistent with the paper’s thesis: the deployable, load-bearing capability is certified Tier-1 control; the quantum column is honestly bounded, with its one near-term niche dated and its misses kept in the record. Nothing in table 22 is on the critical path, and everything in table 21 clears a pre-registered bar. The contrast between the two tables is itself the argument for the tier separation: the Tier-1 column is a list of demonstrated results clearing acceptance bars, while the Tier-3 column is a list of estimates and negatives with a single dated near-term entry. A programme that reported the two together, undifferentiated, would either overclaim the quantum column or underclaim the AI one; reporting them as distinct tiers, each held to the same evidentiary standard, is what keeps the account honest. The full serial-by-serial trace is in Appendix 68, so every entry in all three tables descends to its register gate. Figure 34 shows the Tier-1 metrics against their acceptance bars, and figure 35 shows how the same results discharge the sub-claims of the assurance case.



Tier-1 demonstrated metrics against their acceptance bars (KX-L1-A11/A12/A20, BR-L2-A12, SH-037): each clears its target or dominates its baseline. Frozen register values.

(Schematic.) The assurance case, decomposed to frozen evidence



(Schematic.) The assurance case decomposed to frozen evidence (Appendix 42): the top safety claim rests on five sub-claims, each discharged by a demonstrated result and enforced at runtime by the maturity gate.

The companion series

This paper is the umbrella of the Kronos computational programme; each subsystem it summarises is developed to full depth in a dedicated companion, cross-cited by reserved DOI and official page in §24. Table 23 maps the umbrella sections to the companions that carry them, so a reader can descend from any result here to its complete treatment.

COMPANION	TOPIC	DEVELOPS	
four coupled twins ^[2]	multi-rate co-simulation, conservation gates	§2	
certified CBF filter ^[3]	zero-escape verification, honest PINN negative	§4	
MPC burn control ^[4]	reference-governor MPC in the safe envelope	§4	
neural flux surrogates ^[5]	DeepONet/FNO, manifold abstention	§6	
neural equilibrium ^[6]	sub-ms PINN, sensor-dropout robustness	§6	
calibrated uncertainty ^[7]	conformal prediction, abstention	§6	
runtime assurance ^[8]	signed-ledger provenance, hardware failsafe	§9	
OOD gating ^[9]	extrapolation ledger, epistemic throttling	§6	
safe RL ^[10]	policy exploration inside the certified set	§6	
maturity gate ^[11]	PCMM-to-authority mapping	§2	
diagnostic suite ^[12]	GNN reconstructor, virtual-sensor mesh	§11	
optimisation backbone ^[13]	QP/reduced-basis/adjoint engines	§4	
MLOps ^[14]	versioned lineage, gated promotion	§9	
global sensitivity/UQ ^[15]	40,000-sample Monte-Carlo ranking	§17	
Bayesian OED ^[16]	next-experiment prioritisation	§17	
quantum computing ^[17]	Landau damping, ADAPT-VQE	§16	
quantum optimisation/RL ^[18]	QUBO/QAOA, quantum RL	§16	
quantum ML \	sensing ^[19]	kernel concentration, FT resource ledger	§16
verification-first ^[20]	byte-exact provenance catalog	§9	

The computational-series companions and the umbrella sections they develop. Full DOIs and official pages are in §24.

The relationship is deliberately hierarchical. The umbrella states each result once, with its frozen anchor and its place in the architecture; the companion carries the full method, the ablations, and the extended verification. The **5000**-trajectory CBF study lives in the safety-filter companion ^[3]; the full **129 × 129** PINN reconstruction and the sensor-dropout study live in the equilibrium companion ^[6]; the **40,000**-sample global sensitivity analysis lives in the UQ companion ^[15]. Reading the umbrella tells the reader what the programme is and why it holds together; reading a companion tells the reader exactly how a given piece was built and verified. This paper's contribution is the whole — the certifiable, resource-honest system and method — which no single companion states on its own.

Discussion

WHY THE SEPARATION MATTERS

The value of the architecture is that its tiers are honestly separated and its safety layer is certifiable. Hard safety is unconditional in the surrogate error (Theorem 4.10); performance degrades gracefully; quantum sits where it actually is. The result is a system a regulator and an operator can reason about, and a method other programmes can adopt. The three-clock compute partition, the maturity gate, and the signed ledger are the engineering expression of the same principle: authority follows evidence, and the fastest safety action needs the least trust.

The alternative — trying to build a single learned controller that is both maximally capable and provably safe — runs into a fundamental tension: capability wants a large, expressive, data-hungry model, while provability wants a small, transparent, analysable one. The Kronos design resolves the tension by not asking one object to be both. The performance layer is as large and learned as the operator wishes; the safety layer is small, transparent, and proven; and the two are composed so that the proven layer bounds the learned one unconditionally. This is why the separation is not a compromise but the enabling move: it lets the programme adopt the full power of modern machine learning for performance while keeping the safety argument within the reach of a proof and an audit. A programme that insisted on a single provably-safe learned controller would have to sacrifice capability for analysability; a programme that deployed an unbounded learned controller would have no safety argument at all. The composition escapes both horns.

DESIGN IMPLICATIONS

The certified clamp changes what a learned controller is *allowed* to be. Because safety is unconditional in $\bar{\epsilon}$, the performance layer can be as aggressive, as learned, and as data-hungry as the operator wishes, provided every action passes the projection — which is why a deep-RL policy, an MPC, or a human can share the same actuator without a separate safety case for each. Because the clamp is solver-free and constant-time, it fits inside the control period on deterministic hardware, so the safety guarantee costs no latency budget. And because the two magnets are modelled as distinct devices, the same architecture serves the breeder and the burner without conflating their fields — a load-bearing element of the magnet-integrity case.

There is a second implication for how such a plant is licensed. Because the assurance argument is a live mechanism — the maturity gate enforced call-by-call, the signed ledger recording every decision, the acceptance inequalities pre-registered — a regulator can audit the running system against the same evidence the designer used, rather than against a static safety case that may have drifted from the deployed software. The multiplier λ of the safety QP (§4) gives an interpretable, per-tick measure of how hard the performance layer approached the envelope; the extrapolation ledger records every out-of-distribution call; and the model registry ties every actuation to a validated model version. This is a different licensing posture from “prove the controller is safe once”: it is “prove the boundary holds always, and log everything that approaches it,” which is both more honest about a learned controller and more auditable in operation.

WHAT QUANTUM IS FOR

Quantum sensing is a real, bounded, near-term enhancement of the state estimate: single-shot quench detection with 10^9 margin, and a Heisenberg-limited gain reserved for above-knee error-field mapping. Quantum computation is an offline calibration tier: a validated linear-kinetics encoding that becomes early-fault-tolerant around 2029, a validated alloy-chemistry ansatz, and two honest negatives that tell the programme exactly what *not* to wait for. Nothing on the design critical path waits on quantum hardware, and knowing that with the numbers behind it is itself the result.

RELATION TO THE ASSURANCE AND SAFE-AUTONOMY LITERATURE

The design sits at the intersection of three literatures that rarely meet: control-barrier safety theory ^[39,40,41], learned tokamak control ^[32,33], and structured assurance cases. The CBF literature gives the invariance machinery but is usually demonstrated on robotics, not on a nuclear plant with a hard latency budget and a learned model beneath it; the learned-control literature gives the performance but not the guarantee; the assurance-case literature gives the argument structure but not the runtime enforcement. The Kronos contribution is to fuse the three: a solver-free CBF projection that runs on the hot path, beneath a learned MPC, with an ISS argument (Theorem 4.10) that makes safety unconditional in the model error, and a maturity gate (Appendix 42) that is the runtime expression of the assurance case. The result is that the assurance argument is not a document written after the fact but a live mechanism enforced call-by-call — which is what a regulator of an autonomous plant actually needs.

OPEN PROBLEMS

Three problems are named openly. First, elevating the safety guarantee from the reduced-order envelope to the whole plant requires solving the Hamilton–Jacobi reachability problem 24 on a higher-fidelity model and confirming it in hardware-in-the-loop — a defined but substantial computation (Appendix 45). Second, the flux and equilibrium operators are demonstrated at the frozen design point; parametric generalisation across $(\beta_p, I_p, \text{shape})$ under the same acceptance gates is the one open modelling item, and it follows the same training-and-acceptance recipe rather than a new method. Third, QP feasibility near simultaneous multi-limit saturation is an empirical property of the reference governor on the envelope model, not a theorem; proving it, or constructing a governor with a feasibility certificate, is an open control-theoretic question. None of the three is a gap in the present results; each is the next rung of the maturity ladder (§19).

METHOD TRANSFERABILITY

Although every number here is anchored to the Kronos machines, the method is machine-agnostic. The four elements that carry the contribution — a differentiable predictive twin, calibrated uncertainty with abstention, a solver-free certified safety projection, and a maturity gate binding authority to evidence — assume only a control-affine reduced model with a barrier-encodable safe set and a learned forecast with a computable error bound. Any regulated autonomous plant that meets those assumptions can adopt the pattern: the CBF projection needs only the barrier gradient and the input matrix; the conformal wrapper needs only exchangeable calibration data; the maturity gate needs only pre-registered acceptance inequalities. That the same architecture certifies a spherical-tokamak breeder and a tandem-mirror burner — two confinement concepts with different actuators, different limits, and different magnets (§5.8) — is direct evidence of that transferability within fusion; the control-theoretic core transfers beyond it. The value to the wider field is a template for trustworthy autonomy in safety-critical systems where a learned controller must be allowed to pursue performance without ever being trusted to be safe on its own. Aviation, autonomous driving, grid control, and medical devices all share this structure: a capable, increasingly-learned controller that must operate within a hard safety envelope a regulator can check. The Kronos pattern — a proven, solver-free boundary beneath an arbitrary learned controller, with authority bound to evidence and every decision logged immutably — is a candidate template for all of them. What the fusion setting contributes is the extreme case: a ten-order-of-magnitude timescale span, a learned plant model, and a nuclear-grade consequence, which together force the boundary to be provable, constant-time, and auditable. A pattern that works under those constraints is a strong candidate for the less-extreme settings.

SUMMARY OF THE ARGUMENT

It is worth restating the argument in one place, because the paper’s length can obscure its simplicity. A nuclear-regulated compact fusion generator must be controlled across ten orders of magnitude in time, with every actuation bounded and auditable. A learned controller can be capable but cannot, by itself, be proven safe. The resolution is to compose a large learned performance layer with a small proven safety layer: the performance layer (MPC, or a learned policy) pursues the target, and the safety layer (the control-barrier projection) edits the command minimally to keep the plant in a certified safe set, with an invariance proof (Proposition 4.8) and an input-to-state-stability argument (Theorem 4.10) that make safety unconditional in the model error. The performance layer is made predictive by learned neural operators — a PINN equilibrium at **0.139%** and a flux surrogate at $R^2 = \mathbf{0.998}$ — and honest by calibrated uncertainty with abstention, so it acts on a trusted forecast and declines when the forecast is untrustworthy. Every learned component is admitted to authority only by clearing a pre-registered acceptance inequality, so runtime authority follows verification credibility. Beneath everything sits a deterministic hardware floor that fires in **0.055 μ s** without the AI. Quantum sensing sharpens the state estimate at the margin; quantum computation sits offline with its measured negatives and its one dated niche. And every number is a frozen, reproducible anchor in the de-risking register. That is the whole system: a certified boundary beneath a learned controller, made predictive and honest, gated by evidence, and reported — misses included — with reproducible rigour.

WHAT WOULD CHANGE THE CONCLUSIONS

A useful test of an honest analysis is to say what evidence would overturn it. The safety conclusion would be weakened if the reduced-order envelope model were shown not to bound the full-plant dynamics — which is exactly why the Hamilton–Jacobi successor and hardware-in-the-loop confirmation are on the roadmap. The performance conclusions would be revised if the surrogates failed to generalise across the operating plane — which is why parametric generalisation is the named open modelling item. The quantum conclusions would change if a fault-tolerant machine of the required scale arrived earlier than the vendor roadmaps project, or if an algorithmic advance lowered the resource threshold for nonlinear gyrokinetics — which is why the dated crossings are re-checked annually and stated in machine-independent logical qubits. None of these is hidden; each is a named, testable contingency, and the architecture is built so that its load-bearing guarantee — unconditional safety under the declared disturbance bound — survives all but the first, which is itself the subject of a defined roadmap step.

Limitations

Two honest gates bound the claims. First, and load-bearing: *there is no measured quantum advantage for fusion today*. Every quantum result here is a simulator-scale demonstration, a resource estimate, or a reported negative; the classical baseline is cheaper at every demonstrated scale, and the earliest crossover (linear kinetics) is an offline-calibration target near **2029**, not a control-loop capability. Second, the certified safety guarantee is for the reduced-order operating-envelope model under a declared disturbance bound $\|d\| \leq \bar{d}$, not a whole-plant guarantee: elevating it requires the Hamilton–Jacobi reachability of 24 and hardware-in-the-loop confirmation on the roadmap; QP feasibility near simultaneous multi-limit saturation is an empirical result on the envelope model, not a theorem; and the physics-informed equilibrium network, in the variant intended as a fast *native* solver, reached only **1.4%** relative- L^2 error and missed that sharper target — reported, not reframed. Keeping these negatives in the record is part of the method.

Three further scope statements bound the claims honestly. First, the demonstrated surrogate results are at the frozen design point; parametric generalisation across the operating plane is the open modelling item (§19), and while it follows the same acceptance-gated recipe, it is not yet a frozen result. Second, the detection metrics (GNN reconstruction, fused quench precursor) are demonstrated at twin-synthetic fidelity in a deliberately weak-fault regime; hardware confirmation of the precursor on an instrumented winding is the next rung. Third, the quantum resource estimates are order-of-magnitude and depend on algorithmic and vendor-roadmap assumptions stated with them; the dated crossings are re-checked at each annual review and carry the historical slippage risk of any hardware roadmap. None of these caveats weakens the load-bearing result — unconditional safety under the declared disturbance bound — but stating them is what separates a de-risked claim from an optimistic one.

We also decline to over-caveat. The honest gate is stated once: no measured quantum advantage today, and a reduced-order safety guarantee. Beyond that, the demonstrated Tier-1 results are what they are — zero escapes, a seventh-of-target equilibrium error, a dominant fused precursor — and repeating the caveats at every mention would misrepresent the strength of the frozen evidence. The discipline is to name each limitation clearly and once, then to let the results stand: a de-risked programme is neither one that hides its misses nor one that buries its hits under hedging.

Conclusion

The Kronos computational programme is a certifiable, resource-honest system and method for compact-fusion control. Physics-informed AI is deployable now *because* it is deterministically clamped: an MPC controller acts inside a control-barrier safety projection whose forward invariance we prove (Proposition 4.8) and realise as a solver-free QP, catching **50/100 (5000/5000)** in the companion study) would-be safe-set violations and holding the pressure limit with $h_{\min} = +0.000$. The models beneath it are fast and faithful — a PINN equilibrium at **0.139% RMS ψ** and **2877 $\times$** , a flux surrogate at $R^2 = 0.998$ recovering $Q = 3.076$ as **2.998**, a GNN that survives sensor loss, and a fused quench precursor at AUC **0.992** — carry calibrated uncertainty, and are admitted to authority only through pre-registered acceptance inequalities. The whole campaign runs verification-first on the NVIDIA GPU HPC campaign. Quantum sensing is a near-term diagnostics enhancement; quantum computation is an offline calibration tier with no crossover this decade, and we report the QAOA and QML negatives rather than reframing them. Two distinct magnets — breeder centrepost **16.84 T**, burner plug **26.49 T** — are modelled as such throughout. The contribution is a system and a method, not new plasma physics; each result is a frozen anchor in the de-risking register and heads a family of companion papers in the Kronos 2026 series.

The wider significance is a template for trustworthy autonomy in a safety-critical system. The recurring difficulty with learned control in regulated settings is that capability and provability pull in opposite directions; the resolution demonstrated here is to stop asking a single object to be both, and instead to compose a large learned performance layer with a small proven safety layer so that the proven layer bounds the learned one unconditionally. The mathematics that makes this work — a solver-free control-barrier projection with an analytic invariance proof and an input-to-state-stability argument that renders safety independent of the model error — is not specific to fusion; the frozen anchors are. That the same architecture certifies a spherical-tokamak breeder and a tandem-mirror burner, two confinement concepts with different actuators, limits, and magnets, is the evidence that the method travels. And the quantum column, reported with its measured negatives, is the model of how a regulated programme should treat a promising but immature technology: test each claim, date the one real near-term niche, and keep the misses in the record with the same discipline as the hits. Physics-informed AI is deployable now because it is clamped; quantum sits where it actually is; and both are held to the same standard of frozen, reproducible, auditable evidence.

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. Every headline quantity is a frozen anchor in the KRONOS Physics De-Risking Register ^[1] ([doi:10.5281/zenodo.22133057](https://doi.org/10.5281/zenodo.22133057)) and is regenerable from the archived, uniquely-named figure-generation scripts deposited with this paper; this paper is archived at [doi:10.5281/zenodo.22645702](https://doi.org/10.5281/zenodo.22645702); official version: <https://www.kronosfusionenergy.com/publications/ai-quantum-control.html>. As the umbrella of the computational series, this paper's results are developed in dedicated companion papers, cross-cited here by their reserved DOIs and official pages: the four coupled digital twins ^[2] ([doi:10.5281/zenodo.22645805](https://doi.org/10.5281/zenodo.22645805), <https://www.kronosfusionenergy.com/publications/four-coupled-digital-twins-for-a-compact-spherical-toka.html>); the certified control-barrier safety filter ^[3] ([doi:10.5281/zenodo.22645716](https://doi.org/10.5281/zenodo.22645716), <https://www.kronosfusionenergy.com/publications/control-barrier-safety-clamp.html>); model-predictive burn control inside the safe envelope ^[4] ([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>); neural-operator flux surrogates ^[5] ([doi:10.5281/zenodo.22645803](https://doi.org/10.5281/zenodo.22645803), <https://www.kronosfusionenergy.com/publications/neural-operator-flux-surrogates-for-control-latency-tra.html>); sub-millisecond neural equilibrium reconstruction ^[6] ([doi:10.5281/zenodo.22645801](https://doi.org/10.5281/zenodo.22645801), <https://www.kronosfusionenergy.com/publications/sub-millisecond-neural-equilibrium-reconstruction-robust.html>); calibrated uncertainty, conformal prediction and abstention ^[7] ([doi:10.5281/zenodo.22645813](https://doi.org/10.5281/zenodo.22645813), <https://www.kronosfusionenergy.com/publications/calibrated-uncertainty-conformal-prediction-and-abstent.html>); runtime assurance with signed-ledger provenance ^[8] ([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>); out-of-distribution detection and epistemic gating ^[9] ([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>); safe reinforcement learning from operator feedback ^[10] ([doi:10.5281/zenodo.22645823](https://doi.org/10.5281/zenodo.22645823), <https://www.kronosfusionenergy.com/publications/safe-reinforcement-learning-from-operator-feedback-for.html>); maturity-gated actuation authority ^[11] ([doi:10.5281/zenodo.22645795](https://doi.org/10.5281/zenodo.22645795), <https://www.kronosfusionenergy.com/publications/maturity-gated-actuation-authority.html>); the diagnostic suite and GNN reconstructor ^[12] ([doi:10.5281/zenodo.22645943](https://doi.org/10.5281/zenodo.22645943), <https://www.kronosfusionenergy.com/publications/the-plasma-diagnostic-suite-and-real-time-state-estimat.html>); the real-time optimisation backbone ^[13] ([doi:10.5281/zenodo.22645899](https://doi.org/10.5281/zenodo.22645899), <https://www.kronosfusionenergy.com/publications/a-real-time-optimization-backbone-for-fusion-plasma-con.html>); MLOps for the safety-critical stack ^[14] ([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>); global sensitivity and uncertainty quantification ^[15] ([doi:10.5281/zenodo.22645893](https://doi.org/10.5281/zenodo.22645893), <https://www.kronosfusionenergy.com/publications/global-sensitivity-and-uncertainty-quantification-acros.html>); Bayesian optimal experimental design and multi-fidelity data fusion ^[16] ([doi:10.5281/zenodo.22645895](https://doi.org/10.5281/zenodo.22645895), <https://www.kronosfusionenergy.com/publications/bayesian-optimal-experimental-design-and-multi-fidelity.html>); resource-honest quantum computing against the classical frontier

^[17] ([doi:10.5281/zenodo.22645718](https://doi.org/10.5281/zenodo.22645718), <https://www.kronosfusionenergy.com/publications/quantum-computing-for-fusion.html>); quantum optimisation and reinforcement learning without claimed advantage ^[18] ([doi:10.5281/zenodo.22645903](https://doi.org/10.5281/zenodo.22645903), <https://www.kronosfusionenergy.com/publications/quantum-optimization-and-reinforcement-learning-for-rea.html>); and quantum machine-learning surrogates, entanglement-enhanced magnetometry and fault-tolerant resource estimates ^[19] ([doi:10.5281/zenodo.22645905](https://doi.org/10.5281/zenodo.22645905), <https://www.kronosfusionenergy.com/publications/quantum-machine-learning-surrogates-entanglement-enhanc.html>); with the verification-first provenance catalog ^[20] ([doi:10.5281/zenodo.22645710](https://doi.org/10.5281/zenodo.22645710), <https://www.kronosfusionenergy.com/publications/verification-and-validation-47-code-provenance.html>). No result depends on unpublished or proprietary data.

Scope — what this paper does not claim. The certified safety guarantee is for the reduced-order operating-envelope model under a declared disturbance bound, not a whole-plant guarantee; and quantum computation shows no measured hardware advantage for fusion this decade.

Data availability tocsectionData availability The frozen values behind every figure and table are archived with this deposit ([doi:10.5281/zenodo.22645702](https://doi.org/10.5281/zenodo.22645702)) and trace to the register ^[1] ([doi:10.5281/zenodo.22133057](https://doi.org/10.5281/zenodo.22133057)) by the serial identifiers cited inline. The uniquely-named figure-generation scripts are included in the deposit. Any economic analysis is reported separately and is not part of the public record.

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

section0 Asection Asection.subsection Asection.figure Asection.table

Nomenclature

Table 25 collects the principal symbols. longtablell Principal symbols.\ symbol & meaning \ 2l(Table 25 continued)\ symbol & meaning \ $\mathbf{x}, \mathbf{u}$ & plant state, actuation \ $\mathbf{f}_c, \mathbf{g}_c$ & control-affine drift and input matrix, 7 \ $\mathbf{d}, \bar{\mathbf{d}}$ & disturbance and its bound, 7 \ $\mathbf{h}, \mathcal{C}$ & barrier function, safe set, 16 \ α, γ & extended class- $\mathcal{K}$ function; $\alpha(\mathbf{h}) = \gamma\mathbf{h}$ \ κ_d & robust barrier tightening, 19 \ $\mathbf{L}_f\mathbf{h}, \mathbf{L}_{g_c}\mathbf{h}$ & Lie derivatives of $\mathbf{h}$ \ $\mathbf{Q}, \mathbf{R}$ & MPC tracking and effort weights, 14 \ N & MPC horizon (§4); ensemble/collocation size (context) \ ψ, Δ^* & poloidal flux; Grad–Shafranov operator, 26 \ $\mathcal{G}_\theta, \mathcal{E}_\phi$ & neural flux and equilibrium operators \ $\mathbf{f}_G, \beta_N, q_{95}$ & Greenwald fraction, normalized pressure, edge safety factor \ $\mathbf{P}^f, \mathbf{P}^a, \mathbf{K}$ & forecast/analysis covariance, Kalman gain, 33–33 \ $\delta\mathbf{B}, \gamma_{NV}$ & NV sensitivity, gyromagnetic ratio, 37 \ N_{iter}, N_{phys}, d & Grover iterations, physical qubits, code distance, 43 \ Q, Q_E & breeder fusion gain, burner engineering gain \ B_0 & toroidal field on axis \ D_M & Mahalanobis distance, 36 \ longtable

Table 25 lists the acronyms. longtablell Acronyms.\ acronym & expansion \ 2l(Table 25 continued)\ acronym & expansion \ AD & automatic differentiation \ AIC & Alfvén ion-cyclotron (mode) \ AUC & area under the receiver-operating characteristic \ CBF & control barrier function \ DCLC & drift-cyclotron loss-cone (mode) \ DEC & direct energy conversion \ DFT & density-functional theory \ EnKF & ensemble Kalman filter \ FNO & Fourier neural operator \ FPGA & field-programmable gate array \ FPR / TPR & false-/true-positive rate \ FT & fault-tolerant \ GNN & graph neural network \ HIL & hardware-in-the-loop \ HJ(I) & Hamilton–Jacobi (–Isaacs) \ ISS & input-to-state stability \ KKT & Karush–Kuhn–Tucker \ MHD & magnetohydrodynamics \ MPC & model-predictive control \ NRMSE & normalised root-mean-square error \ NV & nitrogen-vacancy \ NZPV & normal-zone propagation velocity \ OOD & out-of-distribution \ PINN & physics-informed neural network \ POD/ROM & proper-orthogonal decomposition / reduced-order model \ QAOA & quantum approximate optimisation algorithm \ QP & quadratic program \ QPE & quantum phase estimation \ QUBO & quadratic unconstrained binary optimisation \ REBCO & rare-earth barium copper oxide (superconductor) \ RL & reinforcement learning \ SQL & standard quantum limit \ ST & spherical tokamak \ TBR & tritium breeding ratio \ VQE & variational quantum eigensolver \ longtable

The reduced operating-envelope model

The actuator map from physical commands to $(\dot{f}_G, \dot{\beta}_N, \dot{q}_{95})$ is linearised about the operating point. Fuelling rate S_{fuel} drives $\dot{f}_G = (S_{\text{fuel}} - \bar{n}_e/\tau_p)/n_G$ with particle confinement time τ_p ; auxiliary power P_{aux} drives $\dot{\beta}_N$ through the energy balance $dW/dt = P_{\text{heat}} - W/\tau_E$ with $\beta_N \propto W$ at fixed I_p, a, B_0 (10); the current/coil command drives $\dot{q}_{95}$ through $q_{95} \propto a^2 B_0 / (R_0 I_p)$. Collecting these gives the diagonal $g_c = \text{diag}(1/n_G, \partial\beta_N/\partial W \cdot \eta_{\text{heat}}, \partial q_{95}/\partial I_p \cdot \eta_{\text{cd}})$ of 7, where $\eta_{\text{heat}}, \eta_{\text{cd}}$ are the heating and current-drive efficiencies. The diagonality is the structural fact that each actuation channel drives one safe-set coordinate, which is what decouples the multi-barrier clamp into scalar projections (§4); it is a consequence of choosing the state 6 to be the three limit-carrying coordinates rather than an arbitrary basis. The drift f_c collects the passive terms — the density decay $-\bar{n}_e/(\tau_p n_G)$, the energy loss $-W/\tau_E$ mapped to $\dot{\beta}_N$, and the resistive current evolution mapped to $\dot{q}_{95}$ — and is identified from the twin by automatic differentiation at the operating point.

The disturbance $\bar{d}$ aggregates the unmodelled transport (the difference between the reduced energy balance and the full gyrokinetic transport), the coupling residuals of 1–1 bounded by the conservation gates 3–4, and the surrogate error $\bar{e}$ of Theorem 4.10; its bound $\bar{d}$ is declared conservatively and is the scope of the guarantee. Because the barrier gradients are the coordinate axes (Appendix 4.3), the robust tightening $\kappa_{d,i} = \bar{d}_i \|\nabla h_i\| = \bar{d}_i$ reduces to the per-coordinate disturbance bound, so each scalar projection carries exactly the tightening its own channel's uncertainty warrants — a channel with a well-characterised actuator needs little tightening, one with a noisy actuator more.

Standing assumptions

The guarantees of §4–§5 rest on the following assumptions, stated here so the scope is explicit.

[(A1)] *Control-affine reduced model*. The operating point obeys $\dot{x} = f_c(x) + g_c(x)u + d$ with f_c, g_c locally Lipschitz and g_c diagonal on the envelope (Appendix 26).

[(A2)] *Bounded disturbance*. $\|d(t)\| \leq \bar{d}$ with $\bar{d}$ the declared bound aggregating transport, coupling, and estimation error (§3).

[(A3)] *Relative degree one* for the density, β , and kink barriers ($L_{g_c}h_i \neq 0$); the vertical-stability barrier is relative degree two and uses the exponential-CBF extension (Appendix 30).

[(A4)] *Smooth barriers*. Each h_i is continuously differentiable on a neighbourhood of $\mathcal{C}$, with the coordinate-axis gradients of Appendix 43.

[(A5)] *Calibrated estimate*. The state estimate carries a covariance whose empirical coverage clears the 3% conformal gate (SH-027/073), so the estimation-error bound used in the robust tightening is sound.

[(A6)] *Provenance-gated model*. The surrogate in the loop is the validated version, so its error bound $\bar{e}$ (Theorem 4.10) holds (§10).
itemize

Under (A1)–(A6), Proposition 4.8 (forward invariance) and Theorem 4.10 (ISS) hold on the reduced-order operating envelope. Elevating the guarantee to the whole plant relaxes the reduced-model restriction in (A1) and requires the Hamilton–Jacobi successor of Appendix 45 plus hardware-in-the-loop confirmation (§23). The assumptions are all either structural (A1, A3, A4), declared and conservative (A2), or enforced at runtime by an acceptance gate (A5, A6) — none is a hidden idealisation.

Condensed MPC quadratic program

Write the linearised twin about the shadow state as $\mathbf{x}_{k+1} = \mathbf{A} \mathbf{x}_k + \mathbf{B} \mathbf{u}_k$ with $\mathbf{A} = \partial \mathbf{f} / \partial \mathbf{x}$, $\mathbf{B} = \partial \mathbf{f} / \partial \mathbf{u}$ evaluated at $(\mathbf{x}_k, \mathbf{u}_k)$ by reverse-mode automatic differentiation. Unrolling the recursion from $\mathbf{x}_0$,

$$\mathbf{x}_k = \mathbf{A}^k \mathbf{x}_0 + \sum_{j=0}^{k-1} \mathbf{A}^{k-1-j} \mathbf{B} \mathbf{u}_j,$$

and stacking $\mathbf{X} = (\mathbf{x}_1, \dots, \mathbf{x}_N)$, $\mathbf{U} = (\mathbf{u}_0, \dots, \mathbf{u}_{N-1})$ gives the prediction $\mathbf{X} = \mathbf{S}_x \mathbf{x}_0 + \mathbf{S}_u \mathbf{U}$ with

$$\mathbf{S}_x = \begin{bmatrix} \mathbf{A} \\ \mathbf{A}^2 \\ \vdots \\ \mathbf{A}^N \end{bmatrix}, \quad \mathbf{S}_u = \begin{bmatrix} \mathbf{B} & & & \\ \mathbf{A} \mathbf{B} & \mathbf{B} & & \\ \vdots & & \ddots & \\ \mathbf{A}^{N-1} \mathbf{B} & \mathbf{B} & \dots & \mathbf{B} \end{bmatrix},$$

$\mathbf{S}_u$ block-lower-triangular. Substituting into the cost of 14 with $\bar{\mathbf{Q}} = \text{blkdiag}(\mathbf{Q}, \dots, \mathbf{Q}, \mathbf{Q}_N)$ (the terminal weight $\mathbf{Q}_N$ discussed below) and $\bar{\mathbf{R}} = \text{blkdiag}(\mathbf{R}, \dots, \mathbf{R})$ yields the condensed QP 15 with

$$\mathbf{H} = \mathbf{S}_u^\top \bar{\mathbf{Q}} \mathbf{S}_u + \bar{\mathbf{R}}, \quad \mathbf{c} = \mathbf{S}_u^\top \bar{\mathbf{Q}} (\mathbf{S}_x \mathbf{x}_0 - \mathbf{X}^*),$$

where $\mathbf{X}^*$ stacks the target and the state/input path constraints assemble into $\mathbf{G} \mathbf{U} \leq \mathbf{w}$. $\mathbf{H} \succ \mathbf{0}$ because $\bar{\mathbf{R}} \succ \mathbf{0}$, so 15 is a strictly convex QP with a unique minimiser, solvable within the control period by an active-set or operator-splitting method [38].

Terminal ingredients and nominal stability. Choosing the terminal weight $\mathbf{Q}_N$ as the solution $\mathbf{P}$ of the discrete algebraic Riccati equation $\mathbf{P} = \mathbf{A}^\top \mathbf{P} \mathbf{A} - \mathbf{A}^\top \mathbf{P} \mathbf{B} (\mathbf{R} + \mathbf{B}^\top \mathbf{P} \mathbf{B})^{-1} \mathbf{B}^\top \mathbf{P} \mathbf{A} + \mathbf{Q}$ and a terminal set $\mathcal{X}_f$ invariant under the associated LQR feedback makes the optimal cost a Lyapunov function, giving nominal asymptotic stability of the tracking error by the standard argument [36,37]. This is the Lyapunov function V invoked in Theorem 4.10; the reference governor keeps $\mathbf{x}^*$ reachable within the horizon so that the terminal constraint stays feasible. The standard argument runs as follows: the optimal cost $J^*(\mathbf{x})$ is a candidate Lyapunov function; shifting the optimal input sequence by one step and appending the terminal LQR action gives a feasible (sub-optimal) sequence at the next state whose cost is $J^*(\mathbf{x})$ minus the stage cost plus the terminal decrease, and optimality then gives $J^*(\mathbf{x}_{k+1}) \leq J^*(\mathbf{x}_k) - \ell(\mathbf{x}_k, \mathbf{u}_k)$, so J^* decreases along the closed loop and the tracking error is asymptotically stable. The terminal set $\mathcal{X}_f$ must be invariant under the LQR feedback and contained in the state constraints, which the reference governor ensures by shaping the target. This nominal stability is the performance guarantee; the safety guarantee is separate and stronger, holding even where the nominal argument's assumptions (an accurate model) fail. On the hot path the safety projection is the closed form 20, not an iterative solve, so the safety guarantee costs no QP-solve latency even when the performance QP is warm-started across ticks.

Warm starting. Because the operating point moves little between ticks, the previous tick's solution $\mathbf{U}_{k-1}^*$, shifted by one step, is an excellent initial guess for tick k , so the operator-splitting solver typically converges in a handful of iterations [38]. This keeps the performance QP inside the millisecond budget despite the horizon N ; and if a tick's solve fails to converge in time, the safety projection still runs on whatever nominal command is available, so a slow performance solve degrades performance, never safety — the hard/soft split enforced once more, now at the level of solver timing.

KKT derivation of the closed form

For a single active barrier the QP 18 is $\min_{\mathbf{u}} \frac{1}{2} \|\mathbf{u} - \mathbf{u}_{\text{nom}}\|^2$ s.t. $\mathbf{a}^\top \mathbf{u} \geq \mathbf{b}$ with $\mathbf{a} = (\mathbf{L}_{g_c} \mathbf{h})^\top$ and $\mathbf{b} = -\mathbf{L}_{f_c} \mathbf{h} - \boldsymbol{\alpha}(\mathbf{h}) + \boldsymbol{\kappa}_d$. The Lagrangian is $\mathcal{L} = \frac{1}{2} \|\mathbf{u} - \mathbf{u}_{\text{nom}}\|^2 - \lambda(\mathbf{a}^\top \mathbf{u} - \mathbf{b})$, $\lambda \geq 0$. Stationarity gives $\mathbf{u}^* = \mathbf{u}_{\text{nom}} + \lambda \mathbf{a}$. If the unconstrained $\mathbf{u}_{\text{nom}}$ is feasible ($\mathbf{a}^\top \mathbf{u}_{\text{nom}} \geq \mathbf{b}$) then $\lambda = 0$ by complementary slackness. Otherwise the constraint is active, $\mathbf{a}^\top \mathbf{u}^* = \mathbf{b}$, so $\mathbf{a}^\top \mathbf{u}_{\text{nom}} + \lambda \|\mathbf{a}\|^2 = \mathbf{b}$ and $\lambda = (\mathbf{b} - \mathbf{a}^\top \mathbf{u}_{\text{nom}}) / \|\mathbf{a}\|^2 \geq 0$. Combining both cases, $\lambda = [\mathbf{b} - \mathbf{a}^\top \mathbf{u}_{\text{nom}}]_+ / \|\mathbf{a}\|^2$, and substituting $\mathbf{a}, \mathbf{b}$ gives 20. The map $\mathbf{u}_{\text{nom}} \mapsto \mathbf{u}^*$ is the Euclidean projection onto the half-space $\{\mathbf{a}^\top \mathbf{u} \geq \mathbf{b}\}$, hence non-expansive and continuous (Lemma 5.2).

Multiple constraints. For r simultaneously-active barriers the QP is $\min_{\mathbf{u}} \frac{1}{2} \|\mathbf{u} - \mathbf{u}_{\text{nom}}\|^2$ s.t. $\mathbf{A}\mathbf{u} \geq \mathbf{b}$ with $\mathbf{A} \in \mathbb{R}^{r \times m}$. The KKT system is $\mathbf{u}^* = \mathbf{u}_{\text{nom}} + \mathbf{A}^\top \boldsymbol{\lambda}$, $\boldsymbol{\lambda} \geq 0$, $\boldsymbol{\lambda} \odot (\mathbf{A}\mathbf{u}^* - \mathbf{b}) = 0$; on the active set $\mathcal{A}$ the multipliers solve $\mathbf{A}_{\mathcal{A}} \mathbf{A}_{\mathcal{A}}^\top \boldsymbol{\lambda}_{\mathcal{A}} = \mathbf{b}_{\mathcal{A}} - \mathbf{A}_{\mathcal{A}} \mathbf{u}_{\text{nom}}$, a small dense linear solve, and an active-set iteration adds or drops constraints until all multipliers are non-negative and all constraints satisfied. For the Kronos envelope the barrier gradients are the coordinate axes and $\mathbf{g}_c$ is diagonal (§3), so $\mathbf{A}_{\mathcal{A}} \mathbf{A}_{\mathcal{A}}^\top$ is diagonal and the solve decouples into the r scalar projections 21 — which is why the multi-barrier clamp is still solver-free and constant-time. The only case requiring the general active-set solve is two barriers loading the same channel with opposing signs, which the reference governor pre-empts (§4.7). *Complexity.* The scalar projection is $\mathcal{O}(1)$ per barrier, so the composite clamp is $\mathcal{O}(r)$ in the number of barriers — constant for the fixed three-barrier envelope. The general active-set solve, in the rare conflicting case, is $\mathcal{O}(r^3)$ for the dense factorisation, still trivial for $r = 3$; but it is never needed on the hot path because the governor removes the conflict before it arises. This is why the safety projection carries no solver-timeout risk: its worst case is a 3×3 solve, and its normal case is three scalar **max** operations.

Exponential CBF for relative-degree-two barriers

If $L_{g_c}h \equiv 0$ (the input does not appear in $\dot{h}$) the barrier has relative degree two and 17 is vacuous. Define the auxiliary barrier $h_2 = \dot{h} + \gamma_1 h$; then $\dot{h}_2 = L_{f_c}^2 h + L_{g_c} L_{f_c} h u + \gamma_1 \dot{h}$, and enforcing $\dot{h}_2 \geq -\gamma_2 h_2$ recovers a relative-degree-one condition on u with the same closed-form solution. Chaining $h \geq 0$ from $h_2 \geq 0$ requires $h(0) \geq 0$ and $\dot{h}(0) \geq -\gamma_1 h(0)$; the exponential-CBF construction guarantees $h(t) \geq 0$ for all t under these initial conditions ^[40]. Concretely, define the vector $\eta = (h, \dot{h})^\top$; the relative-degree-two dynamics give $\dot{\eta} = F\eta + Gu$ with $F = \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix}$ and the input entering only the second row. Enforcing $\dot{h}_2 \geq -\gamma_2 h_2$ with $h_2 = \dot{h} + \gamma_1 h$ is an affine constraint on u (relative degree one in h_2), so the closed form 20 applies with $h \rightarrow h_2$ and the Lie derivatives $L_{f_c} h_2 = L_{f_c}^2 h + \gamma_1 L_{f_c} h$, $L_{g_c} h_2 = L_{g_c} L_{f_c} h$. The chaining lemma then gives $h(t) \geq 0$ provided the initial conditions $h(0) \geq 0$ and $h_2(0) \geq 0$ hold — the standard exponential-CBF result, applied here to the vertical-stability barrier of the $\kappa = 2.0$ plasma, whose position dynamics are second-order in the input.

Ensemble-Kalman algorithm

The predictive shadow is an ensemble of M members, propagated and updated as follows each assimilation step.

Forecast. Advance each member through the learned twin, $\mathbf{x}_k^{f,(m)} = \mathcal{S}_\theta(\mathbf{x}_{k-1}^{a,(m)}, \mathbf{u}_{k-1}) + \boldsymbol{\eta}^{(m)}$, with model-error draw $\boldsymbol{\eta}^{(m)} \sim \mathcal{N}(\mathbf{0}, \mathbf{Q}_{\text{mod}})$.

Sample statistics. Compute $\bar{\mathbf{x}}^f = \frac{1}{M} \sum_m \mathbf{x}_k^{f,(m)}$ and the covariance $\mathbf{P}^f = \frac{1}{M-1} \sum_m (\mathbf{x}_k^{f,(m)} - \bar{\mathbf{x}}^f)(\cdot)^\top$.

Localise and inflate. Replace $\mathbf{P}^f \leftarrow \rho \odot (\lambda^2 \mathbf{P}^f)$ with a Schur localisation matrix ρ (tapering spurious long-range correlations) and inflation factor $\lambda \gtrsim 1$ (countering under-dispersion at finite M).

Gain. $\mathbf{K} = \mathbf{P}^f \mathbf{H}^\top (\mathbf{H} \mathbf{P}^f \mathbf{H}^\top + \mathbf{R})^{-1}$, 33.

Analysis. Update each member with perturbed observations, $\mathbf{x}_k^{a,(m)} = \mathbf{x}_k^{f,(m)} + \mathbf{K}(\mathbf{y}_k + \boldsymbol{\epsilon}^{(m)} - \mathbf{H} \mathbf{x}_k^{f,(m)})$, $\boldsymbol{\epsilon}^{(m)} \sim \mathcal{N}(\mathbf{0}, \mathbf{R})$, 33; the perturbation keeps the analysis spread consistent with $\mathbf{P}^a = (\mathbf{I} - \mathbf{K} \mathbf{H}) \mathbf{P}^f$, 33. enumerate

The perturbed-observation form is what prevents the classic variance collapse of the deterministic update: without $\boldsymbol{\epsilon}^{(m)}$ the analysis ensemble would systematically under-represent the posterior spread, breaking the calibration gate. The credible intervals emitted from $\mathbf{P}^a$ are accepted only under the coverage gate $|\text{Cov}_{\text{emp}}(\boldsymbol{\alpha}) - \boldsymbol{\alpha}| \leq 3\%$ (SH-027/073), which is a stronger, testable requirement than mere spread consistency. In the linear-Gaussian limit and $M \rightarrow \infty$ the recursion coincides with the classical Kalman filter [56]; at finite M the localisation/inflation pair is the practical price of the low-rank sample covariance. The output $(\bar{\mathbf{x}}^a, \mathbf{P}^a)$ is exactly the calibrated state distribution that the safety layer reads in 34 to decide how far to trust the model.

Localisation and inflation, in detail. At finite ensemble size M the sample covariance $\mathbf{P}^f$ has rank at most $M - 1$ and carries spurious long-range correlations from sampling noise. Localisation applies a distance-decaying taper (a Gaspari–Cohn function) elementwise to $\mathbf{P}^f$, zeroing correlations beyond a physical length scale, which both removes the spurious terms and raises the effective rank. Inflation multiplies the ensemble spread by $\lambda \gtrsim 1$ (or adds a small covariance) to counter the systematic under-dispersion that localisation and model error introduce. The pair is tuned so that the emitted credible intervals pass the coverage gate (SH-027/073); an under-inflated filter produces over-confident intervals that fail the gate, an over-inflated one produces needlessly wide intervals that trigger unnecessary abstention. The tuning is therefore not free: it is pinned by the requirement that the distribution the safety layer reads be honestly calibrated.

PINN training and the Grad–Shafranov discretisation

The elliptic operator Δ^* of 26 is discretised on the 129×129 (R, Z) grid by centred finite differences, with the $1/R \partial_R$ term handled by the conservative form $R \partial_R (R^{-1} \partial_R \psi)$. The PINN uses a Fourier-feature input embedding $\gamma(\mathbf{x}) = [\cos(2\pi \mathbf{B} \mathbf{x}), \sin(2\pi \mathbf{B} \mathbf{x})]$ to mitigate spectral bias, a compact multilayer perceptron, and the composite loss 27 minimised by Adam followed by L-BFGS. The PDE residual is evaluated by automatic differentiation at interior collocation points; the boundary loss enforces the coil-set-consistent boundary flux. The reported metrics (RMS ψ **0.139%**, GS residual **0.32%**, X-point local error **4.39%**, **2877 $\times$** speed-up at **0.87 ms**) are the frozen anchors of KX-L1-A11 ^[6].

Discretisation of Δ^ .* On the uniform (R, Z) grid with spacings h_R, h_Z , the conservative form $R \partial_R (R^{-1} \partial_R \psi) + \partial_Z^2 \psi$ is discretised as

$$(\Delta^* \psi)_{ij} \approx R_i \frac{R_{i+1/2}^{-1} (\psi_{i+1,j} - \psi_{ij}) - R_{i-1/2}^{-1} (\psi_{ij} - \psi_{i-1,j})}{h_R^2} + \frac{\psi_{i,j+1} - 2\psi_{ij} + \psi_{i,j-1}}{h_Z^2},$$

which is second-order accurate and preserves the self-adjoint structure of the operator, so the operator residual is a faithful measure of physical consistency rather than a discretisation artifact. *Error decomposition.* The total reconstruction error decomposes into a data-fit term (dominant in the core), a PDE-residual term (dominant where the operator is stiff), and a boundary term (dominant near the coil set); the **0.139%** RMS is the data-fit floor, the **0.32%** operator residual is the PDE term, and the **4.39%** X-point local error is where the flux gradient is steepest and the second derivatives largest — the expected and well-understood hot spot, which the Fourier-feature embedding mitigates but does not eliminate. The honest negative — a native-solver variant reaching only **1.4%** relative- L^2 — is the case where the data-fit anchor is removed and the network must satisfy the PDE and boundary conditions alone; that it misses the sharper target is reported, not reframed.

Neural-operator architectures

DeepONet. The operator is factorised as $\mathcal{G}_\theta(\mathbf{a})(\mathbf{y}) = \sum_{k=1}^p \mathbf{b}_k(\mathbf{a}) \mathbf{t}_k(\mathbf{y}) + \mathbf{b}_0$, where the branch net $\mathbf{b}_k$ acts on the input function $\mathbf{a}$ sampled at fixed sensor locations $\{\mathbf{x}_j\}$ and the trunk net $\mathbf{t}_k$ acts on the query coordinate $\mathbf{y}$. The universal-approximation theorem for operators guarantees that, for any continuous nonlinear operator $\mathcal{G}$ and tolerance ε , there exist p and network weights such that $|\mathcal{G}(\mathbf{a})(\mathbf{y}) - \mathcal{G}_\theta(\mathbf{a})(\mathbf{y})| < \varepsilon$ uniformly on compact sets ^[44]; the factorisation is what lets one trained model answer the flux at any radius the controller asks for.

Fourier neural operator. Each layer computes $\mathbf{v} \mapsto \sigma(\mathbf{W}\mathbf{v} + (\mathcal{K}\mathbf{v}))$ with the spectral convolution

$$(\mathcal{K}\mathbf{v})(\mathbf{x}) = \mathcal{F}^{-1}(\mathbf{R}_\phi \cdot (\mathcal{F}\mathbf{v}))(\mathbf{x}),$$

where $\mathcal{F}$ is the (discrete) Fourier transform, $\mathbf{R}_\phi$ a learned complex weight applied to the lowest $k_{\max}$ modes (higher modes truncated), and $\mathbf{W}$ a pointwise linear skip. Truncation makes the layer resolution-invariant — the same $\mathbf{R}_\phi$ acts on any grid — and the FFT makes the cost $\mathcal{O}(k_{\max} \log k_{\max})$ per layer, so the operator trains on coarse gyrokinetic runs and evaluates on the finer control mesh without retraining ^[45].

Both architectures are $\mathcal{C}^1$ and are differentiated by automatic differentiation, exposing the Jacobian the controller linearises on in 14; the AD Jacobian is accepted only when it matches finite differences to $< 10^{-3}$ (SH-032). Both are trained against held-out CGYRO cases under the NRMSE gate 28, with a physics-informed penalty enforcing the critical-gradient threshold behaviour ^[42,46]. The demonstrated base case is the degree-three ridge operator of §6 ($R^2 = 0.998$, RMSE 0.088 in Q , anchor recovery $3.076 \rightarrow 2.998$, BR-L2-A12), which sets the acceptance bar the full-profile neural operator must clear ^[5]. The manifold-abstention monitor computes the Mahalanobis distance 36 of each query from the training manifold and flags an out-of-envelope call before the surrogate is trusted, so the controller never extrapolates the operator into the loop.

Training and inference. Both operators are trained by Adam with a physics-informed regulariser and a held-out validation split used to select the mode truncation $k_{\max}$ (FNO) or the number of basis functions p (DeepONet). The training set is the space-filling design of the operating box run through CGYRO/TGYRO (§6), so the surrogate interpolates within the envelope and abstains outside it. Inference is a single forward pass — $\mathcal{O}(p)$ for DeepONet, $\mathcal{O}(k_{\max} \log k_{\max})$ per FNO layer — which is what keeps the transport call inside the medium-loop budget alongside the 0.87 ms equilibrium call. The differentiability the controller relies on is a property of the architecture, not of the training: both are compositions of smooth maps, so the exact Jacobian is available by automatic differentiation at inference time and is checked against finite differences before the operator is admitted (SH-032).

Conformal coverage guarantee

Let $s_1, \dots, s_n$ be nonconformity scores on exchangeable calibration data and s_{n+1} the test score. By exchangeability, the rank of s_{n+1} among $\{s_i\}_{i=1}^{n+1}$ is uniform on $\{1, \dots, n+1\}$, so $\mathbb{P}(s_{n+1} \leq \hat{q}) \geq 1 - \alpha$ for $\hat{q}$ the $\lceil (1 - \alpha)(n+1) \rceil$ -th smallest of the calibration scores. Hence the set 35 has marginal coverage $\geq 1 - \alpha$ for *any* underlying model and any finite n ^[59]. The upper bound $\mathbb{P}(u \in C(\mathbf{x})) \leq 1 - \alpha + 1/(n+1)$ holds when the scores are almost surely distinct, so the coverage is tight to $\mathcal{O}(1/n)$; with a calibration set of a few thousand points the interval is within a fraction of a percent of nominal. The Kronos gate additionally requires empirical *conditional* coverage within 3% (SH-027/073), a stronger, testable bar than marginal coverage alone, achieved by the Mondrian (group-wise) construction of Appendix 49. The nonconformity score for a regression surrogate is the normalised residual $s(\mathbf{x}, u) = |u - \hat{u}_\theta(\mathbf{x})|/\sigma_\theta(\mathbf{x})$, which adapts the interval width to the predicted aleatoric variance of §6, so the interval is wider where the model is genuinely more uncertain — and it is exactly this calibrated interval that the safety gate 34 reads.

Fourier–Hermite encoding of linearised Vlasov–Poisson

The one-dimensional electrostatic Vlasov–Poisson system is $\partial_t \mathbf{f} + \mathbf{v} \partial_x \mathbf{f} - (e/m) \mathbf{E} \partial_v \mathbf{f} = 0$, $\partial_x E = -(e/\epsilon_0) \int \mathbf{f} d\mathbf{v}$. Linearising about a Maxwellian $f_0(v) = (2\pi v_{\text{th}}^2)^{-1/2} e^{-v^2/2v_{\text{th}}^2}$, writing $\mathbf{f} = \mathbf{f}_0 + \mathbf{f}_1$, and taking a single Fourier mode $\mathbf{f}_1, \mathbf{E} \propto e^{ikx}$ gives $\partial_t \mathbf{f}_{1,k} + ikv \mathbf{f}_{1,k} = (e/m) \mathbf{E}_k f'_0(v)$. Expanding the velocity dependence in normalised Hermite functions $\hat{H}_n(v)$ (orthonormal against the Maxwellian weight), $\mathbf{f}_{1,k} = \sum_n \mathbf{C}_{k,n}(t) \hat{H}_n(v)$, and using the Hermite recurrence $v \hat{H}_n = \sqrt{n} \hat{H}_{n-1} + \sqrt{n+1} \hat{H}_{n+1}$ turns the streaming term into a tridiagonal coupling of adjacent Hermite moments,

$$\dot{\mathbf{C}}_{k,n} = -ikv_{\text{th}} (\sqrt{n} \mathbf{C}_{k,n-1} + \sqrt{n+1} \mathbf{C}_{k,n+1}) + \mathbf{S}_n \mathbf{E}_k,$$

with a source $\mathbf{S}_n$ nonzero only for the low-order moments coupled to $\mathbf{f}'_0$. Poisson's law closes the electric-field amplitude, $\mathbf{E}_k = -(e/\epsilon_0 ik) c_0 \mathbf{C}_{k,0}$, in terms of the density moment $\mathbf{C}_{k,0}$. Augmenting the state vector with $\mathbf{E}_k$, the truncated system (retaining N Hermite moments) is

$$\dot{\mathbf{c}} = -i\mathbf{H} \mathbf{c}, \quad \mathbf{H} = \mathbf{H}^\dagger,$$

Hermitian because the streaming coupling is real-symmetric-tridiagonal times $k v_{\text{th}}$ and the field coupling is made anti-symmetric-consistent by the augmentation^[81]. Since $\mathbf{H}$ is Hermitian, $e^{-i\mathbf{H}t}$ is unitary and the state is amplitude-encodable in $\lceil \log_2 N \rceil$ qubits; Trotterising $e^{-i\mathbf{H}t} = \prod_j e^{-i\mathbf{H}_j t/r}$ into local terms realises the evolution on the device. On a 5-qubit statevector simulator ($N = 2^5 = 32$ retained modes) the circuit reproduces the analytic Landau rate to 1–7% across three wavenumbers: at $k\lambda_D = 0.5$, $\gamma_{\text{circuit}} = -0.164\omega_p$ against $\gamma_L = -0.153\omega_p$ (ratio 1.07; exact-fit -0.160), and the dispersion $\omega/\omega_p = 1.4157 - 0.1534i$ reproduces the textbook $1.4156 - 0.1533i$. The imaginary part is the Landau damping rate — a genuinely kinetic, collisionless effect, correctly captured by a finite-dimensional Hermitian truncation because the recurrence transfers spectral weight to ever-higher Hermite moments (phase mixing), which the truncation dissipates in a controlled way. The analytic benchmark is the root of the dielectric function $\epsilon(\omega, k) = 1 + \frac{1}{k^2 \lambda_D^2} [1 + \zeta Z(\zeta)] = 0$ with $\zeta = \omega/(\sqrt{2}k v_{\text{th}})$ and Z the plasma dispersion function; solving it at $k\lambda_D = 0.5$ gives $\omega/\omega_p = 1.4156 - 0.1533i$, whose imaginary part is the Landau rate $\gamma_L = -0.153\omega_p$. That the 5-qubit circuit reproduces both the real and imaginary parts to the third decimal is the validation of the encoding. The Hermite truncation at $N = 2^q$ moments is the resource knob: retaining more moments resolves the phase-mixing cascade further and reduces the truncation error, at a cost of one additional qubit per doubling of N (the amplitude encoding uses $\lceil \log_2 N \rceil$ qubits). This logarithmic qubit cost in the resolved phase-space size is exactly why the linear kinetic operator is an early-fault-tolerant target (11–19 logical qubits for realistic grids) while the nonlinear five-dimensional problem, whose non-unitary structure forbids the amplitude encoding, is not. Under the realistic noise model of §16 the field-energy observable decoheres to the maximally-mixed floor $1/2^5 = 0.031$ and the signature is lost: the algorithm is correct, the present hardware is not. The nonlinear, collisional five-dimensional gyrokinetic problem is non-unitary (the $\mathbf{E} \partial_v \mathbf{f}_1$ term couples modes and the collision operator is dissipative) and requires a linearisation-plus-linear-solver embedding of the HHL class^[69], which is the fault-tolerant-horizon target^[82,83,17].

Fault-tolerant resource formulas

The Grover/Durr-Hoyer minimum-finder over 2^n candidates needs $N_{\text{iter}} = 1.6\sqrt{2^n}$ oracle calls ^[78]; the surface code at distance d needs $N_{\text{phys}} = N_{\text{logical}} \times 2d^2$ physical qubits, with d set by the target logical error rate through $p_L \sim (p/p_{\text{th}})^{(d+1)/2}$ for physical error rate p and threshold $p_{\text{th}} \approx 10^{-2}$ ^[79]. Because $N_{\text{phys}} \propto d^2$ but the *logical* count is d -independent, the roadmap of figure 31 is stated in logical qubits and is robust to the overhead assumptions; only the physical count and the wall-clock carry the d -sensitivity. The catalog recomputation (KX-L3-A6) applied these relations to every workload and reproduced 50/50 rows with zero discrepancy. Table 24 gives worked values.

WORKLOAD	LOGICAL QUBITS	N_{iter}	T -COUNT	WINDOW
linear Vlasov kinetics	11–19	—	2×10^5 – 1.5×10^6	early-FT $\sim 2029 \pm 2$
QAOA design QUBO	23–61	—	5×10^3 – 5.5×10^7	early-FT (classically trivial)
QML kernel (realistic)	37	—	8.5×10^3 /circuit	FT, no advantage
6-D kinetic Hamiltonian	~ 44	—	(~ 19 d runtime)	frontier
alloy exact ground state (QPE)	~ 238	—	(~ 38 d runtime)	frontier
nonlinear 5-D gyrokinetics	80–160	—	10^9 – 10^{13}	post-2035
fleet Grover QUBO ($n = 120$)	156	1.8×10^{18}	1.2×10^{23}	never

Worked fault-tolerant resource estimates from 43 (KX-L3-A6/A7). N_{iter} shown where a Grover formulation applies; T -counts are algorithm-dependent order-of-magnitude estimates.

Two worked lines make the endpoints concrete. For the fleet QUBO at $n = 120$, $N_{\text{iter}} = 1.6\sqrt{2^{120}} = 1.6 \times 2^{60} \approx 1.8 \times 10^{18}$ sequential oracle calls — no foreseeable hardware executes 10^{18} oracle calls inside a control-relevant time, so this is a resource boundary, not a route. For the linear Vlasov operator retaining $N \sim 2^4$ – 2^5 Hermite modes (Appendix 35), the amplitude encoding uses $\lceil \log_2 N \rceil = 4$ – 5 qubits for the state plus ancillae for the qubitised walk ^[70], totalling 11–19 logical qubits and $\mathcal{O}(10^5$ – $10^6)$ non-Clifford gates — an early-fault-tolerant target that public roadmaps place near 2029 ± 2 (KX-L3-A7). The gap between these two lines is the whole story of Tier 3: the one useful near-term niche is linear kinetics; everything that would move the machine (nonlinear turbulence, exact chemistry, fleet search) is frontier or never.

Magic-state throughput. The wall-clock estimates assume a magic-state factory that supplies distilled $|T\rangle$ states fast enough not to stall the computation. A single distillation unit produces one output state per $\mathcal{O}(d)$ surface-code cycles, so meeting an algorithm’s T -rate requires parallel factories, adding to the physical-qubit footprint. This is why the 6-D Hamiltonian and alloy-QPE estimates (~ 19 and ~ 38 days) are runtime-bound rather than qubit-bound: the sequential T -depth, divided by the achievable distillation throughput, sets the calendar. The linear Vlasov operator’s 10^5 – 10^6 T -gates need only a modest factory and so fit an early-fault-tolerant machine; the fleet Grover’s 10^{23} T -gates would need a factory throughput no foreseeable machine provides, which is the quantitative content of “never.”

Numerical schemes and the multi-rate schedule

Each twin is advanced by an implicit scheme at its own step h_i ; the orchestrator reconciles on the macro-step and enforces the conservation gates 3–4, rejecting and re-taking a macro-step at smaller Δt_{macro} on violation. The reduced-order (POD/ROM) fast loop runs at $\geq 100\times$ the full-twin rate (SH-093) to carry the protective path, so the microsecond NV precursor is never gated behind the millisecond transport surrogate (figure 22). The differentiable twin's AD Jacobian is accepted only when it matches a finite-difference reference to $\|J^{\text{AD}} - J^{\text{FD}}\|_{\infty} < 10^{-3}$ (SH-032), the precondition for the on-the-fly linearisation of 14.

Multi-rate stability. A multi-rate scheme is stable only if the coupling does not amplify errors between reconciliation points. Treating the coupling terms c_i of 2 explicitly across the macro-step and the intra-twin dynamics implicitly gives an IMEX (implicit-explicit) scheme whose stability requires the macro-step to resolve the fastest *coupling* timescale, not the fastest intra-twin timescale — which is why the partition by natural timescale is admissible: the electromagnetic dynamics are fast but couple to the thermal domain only through the slow nuclear-heating channel 1. The conservation gates 3–4 are the runtime check that this assumption holds: a macro-step that violates energy or flux conservation is precisely one where a coupling has been under-resolved, and it is rejected and retaken at a smaller Δt_{macro} . This adaptive macro-step is what lets the co-simulation exploit the timescale separation of table 1 without silently accumulating a coupling error that would corrupt the shadow the controller acts on.

ISS proof of the clamped loop

Recall the definition: the interconnection is input-to-state stable (ISS) with respect to a disturbance w if there exist a class- $\mathcal{KL}$ function β and a class- $\mathcal{K}$ function ρ such that $\|x(t) - x^*\| \leq \beta(\|x(0) - x^*\|, t) + \rho(\|w\|_\infty)$; here $w = e$ is the surrogate error. Along the clamped closed loop, $\dot{V} = \nabla V \cdot (f_c + g_c u^* + d)$. Write $u^* = u_{\text{MPC}} + \delta u$ with δu the minimal safety correction 20. Under the nominal MPC Lyapunov decrease $\nabla V \cdot (f_c + g_c u_{\text{MPC}}) \leq -c_1 \|x - x^*\|^2$, and with $\|\delta u\|$ bounded on $\mathcal{C}$ and the surrogate error $\|e\| \leq \bar{e}$ entering through d , the cross terms are bounded by $\|\nabla V\| (\|g_c\| \|\delta u\| + \bar{e}) \leq c_2 \bar{e}$ on the compact sublevel set, giving 25. Since the correction δu is active only near $\partial\mathcal{C}$ and vanishes in the interior, the residual set is $\mathcal{O}(\bar{e})$ and shrinks with the surrogate error, while forward invariance (Proposition 4.8) holds independently of $\bar{e}$. $\square$

Discrete-time counterpart. On the sampled loop with step Δt the barrier condition becomes $h(x_{k+1}) \geq (1 - \gamma \Delta t) h(x_k)$, and the same Lyapunov argument gives a practical-ISS estimate $V_{k+1} \leq (1 - c_1 \Delta t) V_k + c_2 \bar{e} \Delta t + \mathcal{O}(\Delta t^2)$. The sampling residual $\mathcal{O}(\Delta t^2)$ is absorbed into the disturbance bound $\bar{d}$ of 7, so the continuous-time guarantee carries to the sampled controller provided $\gamma \Delta t < 1$ — the condition the real-time scheduler enforces through the latency budget of figure 22. Concretely, with the medium loop at $\Delta t = \mathcal{O}(1 \text{ ms})$ and $\gamma = \mathcal{O}(1\text{--}10) \text{ s}^{-1}$, the product $\gamma \Delta t$ is $10^{-3}\text{--}10^{-2} \ll 1$, so the discrete-time barrier condition is a tight approximation of the continuous one and the sampling residual is negligible against the declared disturbance bound. The protective loop, being combinational hardware rather than a sampled controller, carries no such condition: its guarantee is exact because there is no sampling. The two-tier structure thus confines the sampled-guarantee subtlety to the medium loop, where it is trivially satisfied, and leaves the microsecond protective action free of it entirely.

NV-diamond shot-noise sensitivity

An ensemble of N nitrogen-vacancy (NV) centres is interrogated by a Ramsey sequence: a $\pi/2$ pulse prepares a superposition of the $\mathbf{m}_s = 0$ and $\mathbf{m}_s = \pm 1$ ground-state spin sublevels, the spin precesses for a free-evolution time t accumulating a phase $\phi = \gamma_{\text{NV}} \mathbf{B} t$ in a field $\mathbf{B}$, and a second $\pi/2$ pulse maps the phase to a population read out optically. The population estimate has photon shot noise; propagating that noise to the inferred field gives the single-shot uncertainty $\sigma_B = 1/(\gamma_{\text{NV}} C \sqrt{N} t)$ near the optimal working point, where $C \in (0, 1]$ is the readout contrast/figure of merit that folds in spin-projection noise, photon-collection efficiency and initialisation fidelity. Averaging for a total time T with dead-time-free shots of length $\sim t$ gives $N_{\text{shots}} = T/t$ independent estimates, so the field *sensitivity* (uncertainty per unit bandwidth) is

$$\delta B = \sigma_B \sqrt{t} = \frac{1}{\gamma_{\text{NV}} C \sqrt{N} t} \xrightarrow{t \rightarrow T_2^*} \frac{1}{\gamma_{\text{NV}} C \sqrt{N T_2^* t}},$$

the coherence-limited form used in 37, where the optimal free-evolution time is set by the ensemble dephasing time T_2^* . Substituting the conservative in-winding values $N = 10^{12}$, $T_2^* = 1 \mu s$, $C = 0.03$, and $\gamma_{\text{NV}} = 2\pi \times 28.03$ Hertz gives $B \approx 0.189$ (KX – L1 – A22). The measurement bandwidth is set by the shot period $\tau_{\text{shot}} = 1/(\gamma_{\text{NV}} C \sqrt{N T_2^* t})$, $BW = 1/(2 \tau_{\text{shot}}) = 125$ for $\tau_{\text{shot}} = 4 \mu s$. A no-insulation turn that quenches sheds its screening current $I_{\text{turn}} = 1 kA$, producing at standoff $d = 2 mm$ the azimuthal field step $\Delta B = \mu_0 I_{\text{turn}} / (2\pi d) \approx 100 mT$ of 38; the per-shot signal-to-noise ratio is $\Delta B / (\delta B \sqrt{BW}) \sim 1.5 \times 10^9$, so a single $4 \mu s$ shot resolves the precursor with the lead time set by the millisecond normal-zone crossing rather than the microsecond sensor latency. The $\propto N^{-1/2}$ scaling of 53 is the standard-quantum-limit envelope of figure 19; an entanglement-enhanced (Heisenberg) protocol scales as N^{-1} but only pays above the $\sigma \approx 0.14$ knee of §8.

Explicit QUBO encodings of the three breeder problems

Each control-relevant combinatorial problem is written as a quadratic unconstrained binary optimisation (QUBO), $\min_{\mathbf{x} \in \{0,1\}^n} \mathbf{x}^\top \mathbf{Q} \mathbf{x}$, with hard constraints folded in by quadratic penalties. *Cadence scheduling* assigns n maintenance/duty windows to balance a fleet load ℓ_j ; the number-partition penalty $(\sum_j \ell_j (2x_j - 1))^2$ expands to $Q_{jk} = 4\ell_j \ell_k$ ($j \neq k$), $Q_{jj} = 4\ell_j(\ell_j - \sum_k \ell_k)$. *Blanket layout* places discrete breeding modules on a lattice to maximise a coverage objective subject to one-module-per-slot constraints $\lambda(\sum_{m \in s} x_m - 1)^2$, giving off-diagonal penalties 2λ within a slot. *Coil-current / actuator scheduling* sequences n discrete actuator set-points under rate and simultaneity constraints, with the tracking cost linearised about the operating point contributing the diagonal and the constraint penalties the off-diagonal. Mapping each Q to a depth-three QAOA ansatz^[77] and to a Grover minimum-finder^[78] and applying the surface-code footprint^[79] yields the resource rows of figures 28–28 (KX-L3-A6): **23–61** logical qubits at design scale, **125** at fleet scale, 5×10^3 – 5.5×10^7 T -gates at design scale, and 1.2×10^{23} T -gates (1.8×10^{18} iterations) for the fleet Grover search. Classical simulated annealing matches or beats QAOA on all three to $n = 120$; the exact-to-heuristic crossover is $n = 12$ –**16** (BR-SX-06). The constraint-penalty weights λ are chosen large enough that violating a hard constraint costs more than any feasible objective improvement ($\lambda > \max |\Delta_{\text{obj}}|$ over feasible moves) but not so large that the landscape becomes ill-conditioned. The Ising form follows from $x_i = (1 - Z_i)/2$, giving $H_C = \sum_i h_i Z_i + \sum_{i < j} J_{ij} Z_i Z_j + \text{const}$ with fields and couplings read off from Q ; this is the H_C the QAOA cost unitary $e^{-i\gamma H_C}$ implements. The resource footprints follow from the number of nonzero couplings (setting the circuit depth) and the surface-code distance (setting the physical-qubit multiplier), both recomputed with zero discrepancy across the 50-row ledger.

Quantum-kernel concentration

A quantum-kernel classifier embeds a feature vector $\mathbf{a}$ as a state $|\phi(\mathbf{a})\rangle = U(\mathbf{a})|0\rangle^{\otimes q}$ and classifies on the Gram matrix $K(\mathbf{a}, \mathbf{a}') = |\langle\phi(\mathbf{a})|\phi(\mathbf{a}')\rangle|^2$ ^[80]. For an expressive (approximately Haar-random) q -qubit feature map the off-diagonal kernel entries concentrate about their mean with variance $\text{Var}[K] \sim 2^{-q}$, so distinct inputs become indistinguishable as q grows and the trained classifier loses discriminating power — exponential kernel concentration, the direct analogue of the barren-plateau phenomenon ^[68,66]. This is exactly the measured behaviour of BR-SX-08: the classifier AUC falls from **0.817** at $q = 4$ to **0.691** at $q = 8$ (figure 29). No amount of scaling recovers advantage, and at a realistic **37**-logical-qubit size the classical Gram computation (≤ 0.09 s) beats the quantum evaluation (≥ 48 s), **9–11** orders of magnitude in projection to fault tolerance. The verdict is “never, for this problem,” not “soon.”

Assurance case and regulatory mapping

The maturity gate of §2 is the runtime expression of a structured assurance argument. The top claim — “the plant remains within its licensed safe envelope under all admissible actions” — is decomposed into sub-claims discharged by the evidence of this paper: (i) *unconditional safety*, discharged by Proposition 4.8 and the 0/100 and 0/5000 zero-escape verifications (KX-L1-A13); (ii) *graceful degradation*, discharged by Theorem 4.10 and the ≥ 0.95 precursor-AUC retention under $39\times$ sensor degradation (SH-037); (iii) *estimator trustworthiness*, discharged by the equilibrium (0.139%), flux ($R^2 = 0.998$) and reconstruction (AUC 0.980) metrics and the conformal coverage gate (SH-027/073); (iv) *deterministic fallback*, discharged by the $0.055\mu s$ hardware clamp that fires without the AI; and (v) *provenance and auditability*, discharged by the signed hash-chained ledger and the byte-exact model promotion of §9. Each learned component is admitted to control authority only through the pre-registered inequalities of table 9; the mapping from evidence to authority is explicit, so an auditor can trace any actuation to the gate that permitted it. This is the sense in which the system is *certifiable*: not that it is certified today, but that its guarantees are checkable and its authority is bound to verified credibility rather than to assertion.

Derivation of the operating-envelope barriers

Greenwald barrier. The empirical density limit is $n_G = I_p / (\pi a^2)$ in units of 10^{20} m^{-3} , MA, m ^[27]. With $I_p = 9.66 \text{ MA}$ and $a = 0.48 \text{ m}$, $n_G = 9.66 / (\pi \times 0.48^2) \approx 13.3 \times 10^{20} \text{ m}^{-3}$. The Greenwald fraction is $f_G = \bar{n}_e / n_G$, and the barrier $h_G = f_G^{\max} - f_G$ of 9 has gradient $\nabla h_G = (-1, 0, 0)^\top$ in the state 6, so it couples only to the fuelling channel through $\dot{f}_G$ (Appendix 26).

Troyon barrier. The ideal-MHD β limit is $\beta_{\max}[\%] = \beta_N^T I_p / (a B_0)$ ^[26], which rearranges to the definition of the normalized pressure $\beta_N = \beta[\%] a B_0 / I_p$: the operating β_N is the Troyon coefficient. Hence the safe cap is a cap on β_N directly, $h_\beta = \beta_N^{\max} - \beta_N$ with $\beta_N^{\max} = 3.5$ (KX-L1-A13), $\nabla h_\beta = (0, -1, 0)^\top$, coupling only to the heating channel through the energy balance $dW/dt = P_{\text{heat}} - W/\tau_E$ with $\beta_N \propto W$ at fixed I_p, a, B_0 .

Kink barrier. The external-kink and sawtooth limits require $q_{95} \geq q_{95}^{\min}$; with $q_{95} \propto a^2 B_0 / (R_0 I_p)$ the barrier $h_q = q_{95} - q_{95}^{\min}$ has $\nabla h_q = (0, 0, +1)^\top$ and couples to the current/coil channel. The three coordinate-aligned gradients and the diagonal input matrix g_c are what make the composite CBF-QP separable into the three scalar projections 21; this is a structural consequence of choosing the state 6 to be exactly the three limit-carrying coordinates. The choice is not arbitrary: a different state basis would couple the barriers and force the general active-set solve of Appendix 29 on the hot path, sacrificing the constant-time property. Choosing the state to be the limit coordinates aligns the geometry of the safe set with the actuation structure, and it is that alignment — not any special-case algebra — that keeps the certified clamp solver-free. The caps $f_G^{\max}$, $\beta_N^{\max} = 3.5$, $q_{95}^{\min}$ are declared margins below the computed physical ceilings; declaring them conservatively is what makes the barrier an operational proxy the clamp can hold with a proof, while the physics companion ^[2] sets the ceilings themselves.

ADAPT-VQE operator pool and mapping

The electronic Hamiltonian $\hat{H} = \sum_{pq} h_{pq} a_p^\dagger a_q + \frac{1}{2} \sum_{pqrs} h_{pqrs} a_p^\dagger a_q^\dagger a_r a_s$ is mapped to qubits by the Jordan–Wigner transform $a_p = \frac{1}{2}(X_p + iY_p) \prod_{j < p} Z_j$, turning each fermionic operator into a weighted sum of Pauli strings. The ADAPT-VQE operator pool $\{\hat{A}_k\}$ consists of spin-adapted single ($a_p^\dagger a_q - \text{h.c.}$) and double ($a_p^\dagger a_q^\dagger a_r a_s - \text{h.c.}$) excitations from the reference determinant. At each iteration the algorithm evaluates the energy gradient of every pool operator, $g_k = \partial E / \partial \theta_k|_{\theta_k=0} = \langle \psi | [\hat{H}, \hat{A}_k] | \psi \rangle$, appends the operator with the largest $|g_k|$ as a new factor $e^{\theta_k \hat{A}_k}$, and re-optimises all $\{\theta_k\}$ by VQE; it terminates when $\max_k |g_k|$ falls below a threshold^[71]. Because the variational energy⁴¹ bounds the true ground state from above, the converged energy is a rigorous upper bound, and for the transition-metal-hydride fragments the procedure reaches $< 1.6 \text{ mHa}$ at the ten-qubit active space (figure 26).

Hamilton–Jacobi reachability successor

The whole-plant successor to the barrier certificate computes the backward-reachable tube of the unsafe set. Let $\mathcal{X}_{\text{unsafe}}$ be the complement of the licensed envelope and $g(\mathbf{x})$ its signed distance ($g < 0$ inside unsafe). The value function

$$V(\mathbf{x}, t) = \max_{d(\cdot)} \min_{u(\cdot)} \min_{\tau \in [t, 0]} g(\mathbf{x}(\tau))$$

solves the Hamilton–Jacobi–Isaacs equation 24 backward in time, and its zero-sublevel set $\{\mathbf{x} : V(\mathbf{x}, t) \leq 0\}$ is the set of states from which the disturbance can force the plant into $\mathcal{X}_{\text{unsafe}}$ within $|t|$ despite the best control. Certifying $\mathcal{X}_0$ safe over the horizon is then $\mathcal{X}_0 \cap \{V \leq 0\} = \emptyset$. The CBF of §4 is a provably-conservative inner approximation: any state with $h(\mathbf{x}) \geq 0$ and the barrier condition satisfied is in the CBF-invariant set, which is contained in the true reachable-safe set. The successor solves 24 offline (neural-network value approximation on the reduced model, verified against a grid solver) and uses the CBF as the online, constant-time enforcement of the learned value function’s safe set. Elevating the guarantee from the reduced-order envelope to the whole plant is exactly this replacement of the local barrier by the global value function, plus hardware-in-the-loop confirmation (§19). The computational cost is the reason the CBF runs today and the HJ successor is on the roadmap: solving 24 on a grid scales exponentially in the state dimension, so it is done offline on a reduced model, and its result is verified against the neural-network value approximation before deployment. The online enforcement remains the constant-time CBF projection, now using a barrier derived from the HJ value function rather than the hand-designed operating-envelope barriers. This two-stage structure — expensive offline certification, cheap online enforcement — is standard in reachability-based safe control and is what makes the whole-plant guarantee tractable at all. The neural-network value approximation deserves a caution: a learned value function is only as trustworthy as its verification, so the roadmap couples it with formal verification of the network (bounding its error against a grid solver on the reduced model) before it is allowed to define the online barrier. Until that verification is complete, the hand-designed operating-envelope barriers of §3 — whose invariance is proved analytically — remain the deployed safe set. The successor thus extends the guarantee’s *scope*, from the reduced envelope to the whole plant, without weakening its *rigour*, because the online enforcement stays a proved projection either way.

Latency budget derivation

The shadow lead-time is $\tau_{\text{lead}} = \tau_{\text{advance}} - \tau_{\text{pipeline}}$, the difference between how far ahead the twin advances per wall-clock second and the end-to-end pipeline latency. For the medium loop the pipeline sums the reconstruction ($\sim$ ms), the equilibrium and flux operators (**0.87** ms plus a transport call), the ensemble-Kalman update ($\sim$ ms), the MPC solve ($\sim$ ms, warm-started), and the closed-form clamp (μ s), giving $\tau_{\text{pipeline}} = \mathcal{O}(\mathbf{5\ ms})$; with a reduced-order twin advancing at $\geq \mathbf{100\times}$ real time (SH-093) the advance term dominates and $\tau_{\text{lead}} \geq \mathbf{50\ ms}$ is met with margin. The fast protective loop is separate: the NV precursor (**4 μ s** shot) and the analytic monitors feed the **0.055 μ s** hardware clamp directly, bypassing the medium loop entirely, so the microsecond protective action is never gated behind the millisecond transport surrogate (figure 22). The condition $\gamma\Delta t < \mathbf{1}$ of Appendix 38 bounds the medium-loop step; the protective loop carries no such dependence because it is combinational.

Diagnostic-graph construction

The diagnostic graph $G = (V, E)$ has one node per analyzer and an edge between two analyzers whenever they constrain a common physical quantity (e.g. Thomson and interferometry both constrain n_e ; magnetics and the NV probe both constrain the local field). Node features are the analyzer's normalised measurement vector plus a health flag; edge features encode the shared quantity and a learned coupling weight. The message-passing update 29 runs L layers of neighbourhood aggregation, so information from a healthy analyzer propagates up to L hops to reconstruct a dropped node. Training uses seeded dropout — channels are randomly ablated during training so the network learns to reconstruct from partial graphs — which is what produces the demonstrated dropout robustness (SH-042/043) and the $39\times$ -degradation retention floor (SH-037). The reconstruction and precursor heads share the graph encoder, so the off-normal score 30 and the state estimate are computed in one forward pass, which is what keeps the reconstruction inside the millisecond budget of figure 22. The graph structure also encodes physical priors that a generic dense network would have to learn: connecting only analyzers that share a physical quantity restricts the hypothesis space to physically-plausible reconstructions, which improves sample efficiency and, more importantly, makes the reconstruction interpretable — an auditor can read which neighbours supplied a reconstructed channel. This interpretability is why the graph formulation is preferred over a black-box imputer despite the latter's occasionally lower reconstruction error: in a regulated plant, a reconstruction whose provenance can be traced is worth more than a marginally better one that cannot.

Reduced-order model for the fast loop

The fast protective loop is carried by a projection-based reduced-order model (ROM). Collecting $\mathcal{S}$ snapshots of the full twin state into $\Xi = [\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(\mathcal{S})}]$ and taking the singular value decomposition $\Xi = \mathbf{U} \Sigma \mathbf{W}^\top$, the proper-orthogonal-decomposition basis is the leading r left singular vectors $\mathbf{U}_r$, chosen so the retained energy fraction $\sum_{i \leq r} \sigma_i^2 / \sum_i \sigma_i^2$ exceeds a threshold ^[49]. Projecting the dynamics onto $\mathbf{U}_r$, $\mathbf{x} \approx \mathbf{U}_r \mathbf{a}$, gives the reduced system $\dot{\mathbf{a}} = \mathbf{U}_r^\top \mathbf{f}(\mathbf{U}_r \mathbf{a}, \mathbf{u})$ of dimension $r \ll n$, evaluated at $\geq 100\times$ the full-twin rate. The ROM is admitted only when its error clears the SH-093 bar ($< 5\%$ at $\geq 100\times$ speed); it is what lets the reduced-order twin serve the millisecond protective path while the full neural operators run on the medium loop (figure 22). Because the ROM is a linear projection of the same physics, its Jacobian is consistent with the full twin's, so the CBF projection acting on the reduced state inherits the invariance guarantee of Proposition 4.8 on the reduced envelope. *Error bound.* The POD projection error is controlled by the discarded singular values: $\|\mathbf{x} - \mathbf{U}_r \mathbf{U}_r^\top \mathbf{x}\|^2 \leq \sum_{i > r} \sigma_i^2$ over the snapshot set, so choosing r to capture $\geq 99.5\%$ of the energy bounds the reconstruction error a priori. The residual is folded, like the sampling and integration residuals, into the disturbance bound $\bar{\mathbf{d}}$, so the reduced-order clamp remains sound; the SH-093 acceptance test ($< 5\%$ at $\geq 100\times$ speed) is the runtime confirmation that the chosen r meets the bound in practice.

Adaptive conformal calibration

The marginal coverage of Appendix 34 can be sharpened to *conditional* coverage by Mondrian (group-wise) conformalisation: partitioning the input space into regimes $\{\mathcal{R}_g\}$ and calibrating a separate quantile $\hat{q}_g$ per regime gives $\mathbb{P}(u \in C(\mathbf{x}) \mid \mathbf{x} \in \mathcal{R}_g) \geq 1 - \alpha$ within each regime, which is what the SH-027/073 gate tests ($|\text{Cov}_{\text{emp}}(\alpha) - \alpha| \leq 3\%$ conditionally, not merely marginally). Under distribution drift, adaptive conformal inference updates the quantile online, $\hat{q}_{t+1} = \hat{q}_t + \eta(\alpha - \mathbf{1}[u_t \notin C_t])$, which provably tracks the target coverage in the long run regardless of the drift: the update raises the threshold when a miss occurs and lowers it when the set covers, so the long-run miss frequency converges to α for any sequence, adversarial included. This matters for a first-of-a-kind machine whose operating distribution shifts as it commissions — the calibration adapts online rather than assuming a fixed distribution, so the coverage guarantee survives the very non-stationarity a new plant exhibits. The three epistemic signals of §6 — conformal set, Mahalanobis distance, predictive variance — are combined by a logical OR: any one firing hands authority to the deterministic floor, so the gate is conservative by construction.

MLOps and provenance

Every model is a versioned artifact carrying its training data hash, its configuration, and its validator. Promotion to a higher authority rung requires passing a two-tier validation gate: Tier-1 byte-exact reproduction of the frozen reference result, and Tier-2 tolerance reproduction against an independent library. A model registry records what is deployed where, at which version, with which validator, so any actuation can be traced to the exact model and data that produced it. Provenance is load-bearing for safety, not paperwork: in the campaign, control-barrier invariance held in **10/10** provenance-gated trials against **1/10** ungated — the difference being that a provenance-gated model is guaranteed to be the validated version, while an ungated one may silently be a stale or mis-configured checkpoint. Every runtime decision — inputs, active barriers, model versions, and the projected command — is written to a signed, hash-chained ledger ^[8], so the audit trail is tamper-evident. No result in this paper depends on unpublished or proprietary data, and each is regenerable from the archived scripts and the register. The MLOps discipline is what makes “verification-first” operational rather than aspirational: a model cannot be deployed without a registry entry recording its data hash, configuration, validator, and reproduction tier, and the registry is queried at deployment to enforce the maturity rung. In effect the acceptance inequalities of table 9 are compiled into the deployment pipeline, so a model that has not cleared its bar cannot be promoted to command authority — the gate is mechanical, not procedural. This is the difference between a process that documents intentions and one that enforces them; the former can drift, the latter cannot without leaving a trace in the signed ledger.

Troyon limit from the ideal-MHD energy principle

The Troyon scaling of 10 follows from the ideal-MHD energy principle. Stability requires the potential-energy functional $\delta W[\xi] \geq 0$ for all admissible displacements ξ ; writing $\delta W = \delta W_{\text{fluid}} + \delta W_{\text{vacuum}}$ and minimising over ξ at fixed pressure and current profiles yields a marginal-stability boundary that, for a fixed shape and profile family, scales the maximum stable β linearly with the normalized current $I_p / (aB_0)$ ^[26]. The proportionality constant is the Troyon coefficient β_N^T , and the normalized pressure $\beta_N = \beta[\%] aB_0 / I_p$ is defined so that the operating point's β_N is directly comparable to β_N^T . For the negative-triangularity breeder the accessible β_N is set by the ballooning and external-kink boundaries computed in the stability companion ^[2]; the control cap $\beta_N^{\text{max}} = 3.5$ (KX-L1-A13) is declared conservatively below that computed ceiling, so the clamp holds an operating margin, not the physical limit itself. This is why the barrier $h_\beta = \beta_N^{\text{max}} - \beta_N$ of 11 is an *operational* proxy: the physics sets the ceiling, the control sets a margin below it, and the clamp proves the margin is never crossed.

Particle-filter alternative to the EnKF

Where the state distribution is strongly non-Gaussian — for instance near a bifurcation in the operating point — the ensemble-Kalman update of Appendix 31, which is exact only for Gaussian statistics, can be replaced by a particle filter. Each particle $\mathbf{x}^{(m)}$ carries a weight $w^{(m)}$ updated by the likelihood $w^{(m)} \propto w^{(m)} p(\mathbf{y} | \mathbf{x}^{(m)})$, and the posterior is the weighted empirical distribution $\hat{p}(\mathbf{x}) = \sum_m w^{(m)} \delta(\mathbf{x} - \mathbf{x}^{(m)})$; systematic resampling when the effective sample size $N_{\text{eff}} = 1 / \sum_m (w^{(m)})^2$ falls below a threshold prevents weight degeneracy. The particle filter is exact as the particle count grows but scales poorly in state dimension, so the design uses the EnKF for the full state and reserves the particle filter for low-dimensional, strongly non-Gaussian sub-problems (e.g. a disruption-precursor mode). Either estimator feeds the same calibrated distribution into the gate 34, and either is accepted only under the conformal coverage bar (SH-027/073); the choice is a fidelity-versus-cost trade, not a safety one, because the guarantee of Proposition 4.8 is independent of the estimator. The practical rule is to use the EnKF for the full, moderately-Gaussian state and to switch to a particle filter only for a low-dimensional sub-state near a bifurcation, where the posterior becomes multimodal and the Gaussian assumption of the EnKF would misrepresent the uncertainty. Because both estimators are wrapped by the same conformal calibration, the switch does not change the coverage guarantee the safety layer relies on — it only changes how faithfully the interior of the distribution is represented. This is the general pattern of the design: the safety guarantee is invariant to the choices made for fidelity, so those choices can be made on cost and accuracy grounds without a separate safety argument for each.

The safe reinforcement-learning objective

The safe-RL layer of §6 learns a policy π_θ that proposes $\mathbf{u}_{\text{nom}} = \pi_\theta(\mathbf{x})$ and executes the projected action $\mathbf{u}^* = \text{proj}_{\mathcal{C}}(\mathbf{u}_{\text{nom}})$. The objective is the expected discounted return under the projected dynamics,

$$J(\theta) = \mathbb{E}_{\pi_\theta} \left[\sum_{t \geq 0} \gamma_{\text{RL}}^t r(\mathbf{x}_t, \mathbf{u}_t^*) \right],$$

and the policy gradient is $\nabla_\theta J = \mathbb{E}[\nabla_\theta \log \pi_\theta(\mathbf{u}_{\text{nom}} | \mathbf{x}) \mathbf{A}(\mathbf{x}, \mathbf{u}^*)]$ with advantage $\mathbf{A}$. Because the projection is differentiable almost everywhere (Appendix 29), the gradient can also flow through $\mathbf{u}^*$ into the policy, so the agent learns to propose actions that need little correction — it internalises the safe set rather than fighting the clamp. The key property is that the safe set is *never* left during training: every executed action is projected, and Proposition 4.8 holds for each, so exploration is confined to $\mathcal{C}$ by construction. This is what makes reinforcement learning admissible in a nuclear loop at all: the exploration that RL requires is exactly what a safety-critical plant cannot tolerate unbounded, and the projection bounds it without bounding the learning.

Worked control-cycle timing

The end-to-end control cycle of §7 fits the medium-loop budget as follows. The GNN reconstruction is a single forward pass ($\sim$ ms); the PINN equilibrium is **0.87 ms** (KX-L1-A11); the flux surrogate is a microsecond-scale polynomial or a single neural-operator pass; the ensemble-Kalman update is a small dense solve ($\sim$ ms); the condensed MPC QP is warm-started from the previous tick and solved by an operator-splitting method in $\sim$ ms; and the CBF projection is the closed form 20, evaluated in microseconds. The sum is $\tau_{\text{pipeline}} = \mathcal{O}(5 \text{ ms})$, well inside the $\geq 50 \text{ ms}$ shadow lead the reduced-order twin provides (SH-093), so the controller acts on a forecast that leads the plant. The protective loop runs beneath all of this: the NV precursor shot ($4 \mu\text{s}$) and the analytic monitors feed the **0.055 μs** hardware clamp directly, on the deterministic FPGA tier, never contending with the medium loop. This separation — a millisecond performance loop and a microsecond protective loop, joined only downward through the data diode — is why the safety action is never gated behind the learned layers.

Bayesian optimal experimental design

The de-risking programme prioritises the next experiment by expected information gain. For a candidate experiment ξ with outcome y and parameters θ , the expected information gain is the mutual information

$$\text{EIG}(\xi) = \mathbb{E}_{y|\xi} [\text{KL}(p(\theta | y, \xi) \| p(\theta))],$$

the expected Kullback–Leibler divergence from prior to posterior. Choosing $\xi^* = \arg \max_{\xi} \text{EIG}(\xi)$ selects the experiment that most reduces uncertainty in the frozen anchors, across reduced-order, GPU-HPC and hardware data sources with their different costs and fidelities ^[16]. Multi-fidelity data fusion combines the sources by weighting each by its calibrated uncertainty, so a cheap low-fidelity run informs the prior for an expensive high-fidelity one. For the control stack this machinery answers a practical question — which measurement would most tighten the error bar on the safety-relevant anchors — and it is the formal complement to the global sensitivity analysis of §17, which identifies *which* anchors matter, where the OED identifies *which experiment* best resolves them.

Cross-machine transfer learning

A disruption/precursor model pre-trained on a source domain $\mathcal{D}_S$ (TCV, DIII-D, NSTX) is transferred to the target $\mathcal{D}_T$ (the compact-ST breeder) by fine-tuning the shared representation while re-weighting for the covariate shift. Writing the target risk as $R_T(\theta) \leq R_S(\theta) + d(\mathcal{D}_S, \mathcal{D}_T) + \lambda^*$ with d a domain-divergence term and λ^* the joint-optimal error, the transfer is useful when the domain divergence is small relative to the target data scarcity — which holds because the precursor physics is shared across devices even when the geometry differs ^[12]. The residual domain shift between conventional aspect ratio and $A = 2.5$ is quantified as an added epistemic-uncertainty term that inflates the predictive variance in the gate 34, so the transferred model is admitted only at the authority its transfer uncertainty warrants and earns higher authority as target data accumulate. This is how a first-of-a-kind machine carries a credible protective model from its first shots without over-trusting it (§11).

Worked NV multiplexing budget

The 10^9 per-shot margin of §8 is spent as follows. Reading M winding locations from one ensemble in a round-robin divides the per-location integration time by M , raising δB by $\sqrt{M}$ per 37. Suppose the winding needs $M = 10^3$ monitored locations for full coverage; then the per-location sensitivity degrades from $0.189 \text{ pT Hz}^{-1/2}$ to $0.189\sqrt{10^3} \approx 6 \text{ pT Hz}^{-1/2}$, and the per-shot SNR on the 100 mT quench step falls from $\sim 1.5 \times 10^9$ to $\sim 5 \times 10^7$ — still seven orders of magnitude above the detection floor. Radiation degradation is multiplicative on top: the SH-037 result tolerates a further $39\times$ loss of the NV channel while holding precursor AUC ≥ 0.95 , which corresponds to a degraded per-shot SNR of $\sim 1.3 \times 10^6$, comfortably single-shot. The multiplexing-times-degradation budget is thus $10^3 \times 39 \approx 4 \times 10^4$, against a 10^9 margin: the picotesla sensitivity is not needed to see a millitesla step, but it is exactly the reserve that lets one ensemble cover the whole winding and survive the neutron field.

Error budget for the Landau circuit

The 1–7% agreement of the 5-qubit Landau circuit (§16) has three error sources. *Hermite truncation*: retaining $N = 32$ moments dissipates the phase-mixing cascade at the highest moment, contributing an error that falls with N ; at $N = 32$ it is sub-percent for $k\lambda_D = 0.5$. *Trotter error*: splitting $e^{-iHt} = (\prod_j e^{-iH_j t/r})^r + \mathcal{O}(t^2/r)$ contributes an error controlled by the Trotter number r ; the reported rate uses r large enough that the Trotter error is below the truncation error. *Fit error*: extracting the damping rate from the field-energy decay by a least-squares fit contributes the remaining spread across the three wavenumbers. The ideal-circuit total is 1–7% (largest at the most weakly-damped wavenumber, where the fit is hardest); the exact-evolution fit (-0.160 at $k\lambda_D = 0.5$) brackets the circuit (-0.164) and the analytic value (-0.153), confirming the discrepancy is dominated by truncation and fit rather than by a circuit-construction error. Under the noise model (single-qubit 10^{-3} , two-qubit 10^{-2} , readout 10^{-2}) the field-energy observable decoheres to $1/2^5 = 0.031$ and the rate cannot be extracted — the pre-registered hardware negative.

Detection metrics

For completeness we define the detection metrics used for the quench precursor and the reconstruction. For a threshold τ on a detector score, the true-positive rate is $\text{TPR}(\tau) = \mathbb{P}(\text{score} > \tau \mid \text{fault})$ and the false-positive rate $\text{FPR}(\tau) = \mathbb{P}(\text{score} > \tau \mid \text{nominal})$; sweeping τ traces the receiver-operating-characteristic (ROC) curve $(\text{FPR}(\tau), \text{TPR}(\tau))$, and the area under it is $\text{AUC} = \int_0^1 \text{TPR} \, d(\text{FPR})$, equal to the probability that a random fault sample scores above a random nominal sample. Under a Gaussian shift model — nominal scores $\mathcal{N}(0, 1)$, fault scores $\mathcal{N}(d', 1)$ — the ROC has $\text{AUC} = \Phi(d'/\sqrt{2})$, the relation 40 that reproduces the frozen anchors. A protection trip lives in the low-FPR corner, so the operating points at 1% and 0.1% FPR (Table 14) matter more than the AUC alone; the fused detector's advantage is largest exactly there, which is why the quadrature fusion 39 is worth its complexity.

Conformal versus Bayesian uncertainty

The uncertainty machinery combines a Bayesian-style ensemble (Appendix 31, §6) with a distribution-free conformal wrapper (Appendix 34), and it is worth saying why both are used. The ensemble/heteroscedastic head produces a predictive distribution that is informative — it says *where* the uncertainty is — but its calibration depends on the model being correct, which a learned surrogate is not guaranteed to be. The conformal wrapper produces a coverage guarantee that is *distribution-free* — it holds regardless of the model — but it is only as sharp as the nonconformity score it is given. Composing them takes the informativeness of the Bayesian estimate and the guarantee of the conformal wrapper: the conformal set is built on the ensemble's normalised residual, so it is both calibrated (conformal) and adaptive to where the model is uncertain (Bayesian). This is why the acceptance gate tests the *conformal* coverage (SH-027/073) rather than the raw ensemble spread: the guarantee the safety layer relies on must not depend on the surrogate being right.

Verification of the reference codes

Each backbone code carries verification evidence before any surrogate is trained against it. The gyrokinetic solvers (CGYRO, GENE) are verified by convergence under refinement of the velocity-space and radial grids and by reproduction of standard linear-benchmark growth rates; the byte-match to the reference growth rate $\gamma = +0.318$ (BR-L2-A18) is the frozen instance. The free-boundary equilibrium solver (FreeGS/FreeGSNKE) is verified by the method of manufactured solutions — prescribing $p'(\psi)$ and $\mathbf{F}F'(\psi)$ so the exact ψ is known — and by grid-refinement convergence of the Grad–Shafranov residual, which is what makes the PINN's 0.32% operator residual a meaningful comparison. The Monte-Carlo neutronics (OpenMC) is verified by analytic-slab and criticality benchmarks and run to statistical convergence with reported variances. The materials stack (DFT with machine-learned interatomic potentials) is verified against higher-level quantum-chemistry references on the transition-metal fragments, the same fragments the ADAPT-VQE check targets. Verification precedes validation, which precedes independent reproduction (§10); only a code that clears all three anchors a surrogate. The ordering matters: a code that solves its equations correctly (verified) but describes the wrong physics (not validated) would anchor a surrogate to a wrong answer, and a code that is both verified and validated but not independently reproducible could not be trusted to have been run correctly on the specific case that froze the anchor. Requiring all three — and stamping the result in the register with a serial — is what lets every headline number in this paper be traced from the surrogate that uses it, through the acceptance gate, to the reference code and the reproduction that froze it. This chain of custody is the concrete meaning of “verification-first”: the burden is on each number to demonstrate its provenance, not on the reader to assume it.

The disruption-precursor library

The precursor detector is trained against a curated library of precursor windows drawn from the source machines and the twin, with strict leakage-free splits: windows from the same pulse never appear in both training and test, and the split is by pulse rather than by window, so temporal autocorrelation cannot inflate the reported metrics. The library is versioned and provenance-stamped like any other training artifact, so the reported AUC (**0.992** fused, KX-L1-A20) is reproducible from the frozen library. The weak-fault regime of the demonstration — per-sample SNR set **10–30×** below the design-basis fault — is chosen deliberately so the detector is stressed where it is hard to discriminate; at the design-basis fault every channel saturates and the discrimination is trivial, so reporting the weak-fault metrics is the conservative, informative choice.

Convergence of the Monte-Carlo verification

The zero-escape result is a Monte-Carlo estimate, so its statistical strength is worth stating. Over N trajectories with random initial conditions and disturbances, observing **0** escapes gives, by the rule of three, a **95%** upper confidence bound on the true escape probability of $\approx 3/N$: 3×10^{-2} for the umbrella $N = 100$ study and 6×10^{-4} for the companion $N = 5000$ study ^[3]. The larger study therefore tightens the empirical bound by two orders of magnitude, complementing — not replacing — the analytic guarantee of Proposition 4.8, which holds for *all* admissible disturbances within the declared bound, not just the sampled ones. The Monte-Carlo and the proof are two independent lines of evidence: the proof establishes that no escape is possible under the stated assumptions, and the Monte-Carlo confirms the implementation realises the proof. The worst-case adversarial trajectory (filtered $h_{\min} = +1.8 \times 10^{-15}$) is the third line: it drives the disturbance to the boundary of the declared bound and confirms the filtered state rides the safe-set boundary without crossing, exactly as the analytic margin predicts.

Burner actuator models

The burner's actuators map to its safe-set coordinates as the breeder's do. Electron-cyclotron heating sustains the plug thermal barrier that decouples the plug electrons from the central cell, raising the plug potential and hence the plug-potential margin $\hbar_\phi$; sloshing-ion injection shapes the ambipolar potential by populating a non-Maxwellian tail, tuning the ion-confining potential; and the direct-energy-converter grid bias dispatches the charged-particle power, which through the power balance affects the central-cell β_c . As for the breeder, the map is linearised about the operating point to give a diagonal input matrix, so the burner clamp decouples into scalar projections over $(\phi_i, \beta_c, \text{microstability})$ with the same closed form 20. The microstability barriers (DCLC and Alfvén ion-cyclotron) are relative-degree-two in the distribution shape and use the exponential-CBF extension of Appendix 30.

Comparison of quantum-simulation approaches

Three approaches to the kinetic problem trade off differently. *Trotterised Hamiltonian simulation* (used here) splits e^{-iHt} into local factors; it is simple, has a clear error budget (Appendix 58), and is the natural fit for the sparse, Hermitian linear Vlasov generator. *Qubitisation* ^[70] achieves optimal scaling in the simulation time and the Hamiltonian norm by encoding H in a block of a larger unitary; it is more efficient asymptotically and is the method assumed in the linear-Vlasov resource estimate (11–19 logical qubits). *Variational* approaches (VQE/ADAPT-VQE) target ground states rather than dynamics and are the right tool for the alloy-chemistry problem, not the kinetic one. The choice per problem is dictated by whether the target is dynamics (Trotter/qubitisation) or a ground state (variational), and by whether fault tolerance is assumed (qubitisation) or not (near-term variational); the resource ledger of Table 24 uses the appropriate method for each row.

The signed-ledger data structure

The runtime ledger is an append-only chain of records. Each record $\mathbf{r}_k$ contains the timestamp (from the PTP/White-Rabbit timebase), the state estimate and its covariance, the active barriers and their multipliers, the model versions in the loop, the nominal and projected commands, and the SHA-256 hash $H(\mathbf{r}_{k-1})$ of the previous record; the record is signed with the controller's key. Because each hash commits to the entire prior chain, altering any past record invalidates every subsequent hash, so an auditor verifying the chain from genesis detects any tampering. The ledger is what makes the multiplier λ of §4 an *auditable* safety price rather than a transient runtime value: every time the clamp intervened, and how hard, is recorded immutably. Combined with the provenance stamps on the training data and the two-tier reproduction gate, the ledger closes the audit loop from datum to model to actuation to log.

Design decisions and their rationale

It is useful to collect the principal design decisions and the reason each was made, because a review of a whole programme risks losing the logic beneath the components.

A boundary, not a behaviour. Safety is enforced by a proven projection rather than trained into the controller, because an empirical safety claim cannot be established across the whole operating space, whereas an analytic invariance guarantee can (§4).

A closed form, not a solver. The projection is the closed form 20 rather than an online QP solve, because the protective action must be bounded in latency on deterministic hardware (§5).

Coordinate-aligned barriers. The state is chosen to be the three limit-carrying coordinates so the clamp decouples into scalar projections, keeping it constant-time (Appendix 43).

Calibrated distributions, not point estimates. The estimator emits a covariance so the clamp can size its robustness margin and the gate can abstain, which requires the conformal coverage guarantee (§6).

A three-clock partition. Compute is split by timescale and security domain, joined by a data diode, so the microsecond protective loop never contends with, and cannot be reached by, the training fleet (§9).

Authority follows evidence. The maturity gate binds runtime authority to pre-registered acceptance inequalities, making the path to autonomy a testable sequence and the deployment reversible (§19).

Frozen, reproducible numbers. Every quantity is a stamped register anchor regenerable from archived scripts, so the paper is a set of reproducible results, not assertions (§10).

Quantum reported honestly. The quantum tiers are placed where the evidence puts them, with measured negatives kept in the record and the one near-term niche dated in machine-independent units (§16). itemize

Each decision follows from the same root requirement: an autonomous nuclear controller must be one a regulator and an operator can reason about, which means its guarantees must be checkable, its authority justified, and its claims reproducible. The components are established; the discipline of assembling them under that requirement is the contribution.

Register-serial glossary

Table 68 lists every de-risking-register serial cited in this paper, its metric, and its status, so each headline number can be traced to its gate. “Demonstrated” denotes a frozen result reproduced in the campaign; “target” a pre-registered bar; “estimate” a resource calculation; “negative” a reported non-result.

longtablelll Register serials cited in this paper. \ serial & metric & status \ 3l(Table 68 continued)\ serial & metric & status \ KX-L1-A11 & PINN equilibrium RMS ψ **0.139%**, **2877 $\times$** , GS residual **0.32%** & demonstrated \ KX-L1-A12 & GNN reconstruction AUC **0.980**, NRMSE **0.396** & demonstrated \ KX-L1-A13 & CBF clamp **0/100 (0/5000)** escapes, $\beta_N = \mathbf{3.500}$, $h_{\min} = +\mathbf{0.000}$ & demonstrated \ KX-L1-A15 & two-magnet load lines: centrepost ΔT **49.5** K, plug **37.1** K & demonstrated \ KX-L1-A20 & fused precursor AUC **0.992**; TPR **0.857/0.597** @ **1%/0.1%** FPR & demonstrated \ KX-L1-A22 & NV magnetometer **0.189** $\text{pT Hz}^{-1/2}$, **125** kHz, SNR $\sim \mathbf{10^9}$ & demonstrated \ BR-L1-A11 & TGYRO confinement $\tau_E = \mathbf{0.410}$ s, $H_{98} \approx \mathbf{0.94}$ & demonstrated \ BR-L2-A12 & flux surrogate $R^2 = \mathbf{0.998}$, RMSE **0.088**, anchor **3.076** $\rightarrow$ **2.998** & demonstrated \ BR-L2-A18 & CGYRO linear growth rate $\gamma = +\mathbf{0.318}$ (byte-match) & demonstrated \ SH-026 & flux-operator NRMSE $< \mathbf{5\%}$ vs held-out CGYRO & target \ SH-031 & equilibrium-operator RMS $\psi < \mathbf{0.5\%}$ at kHz & target \ SH-032 & AD-vs-FD Jacobian $< \mathbf{10^{-3}}$ & target \ SH-037 & precursor AUC $\geq \mathbf{0.95}$ with NV degraded **39 $\times$** & demonstrated \ SH-042/043 & seeded-dropout bounded state error & target \ SH-093 & reduced-order fidelity $< \mathbf{5\%}$ at $\geq \mathbf{100\%}$ & target \ SH-096 & energy/flux conserved across couplings & target \ SH-027/073 & ensemble coverage $|\text{Cov} - \alpha| \leq \mathbf{3\%}$ & target \ SH-102 & three-clock latency budget & design allocation \ KX-L3-A1 & 6-D kinetic Hamiltonian sim $\sim \mathbf{44}$ logical qubits, $\sim \mathbf{19}$ d & estimate \ KX-L3-A2 & alloy QPE $\sim \mathbf{238}$ logical qubits, $\sim \mathbf{38}$ d & estimate \ KX-L3-A6 & FTQC resource ledger, **50/50** rows, **0** fail & demonstrated \ KX-L3-A7 & dated roadmap: early-FT $\sim \mathbf{2029 \pm 2}$; 5-D GK post-**2035**; fleet Grover never & estimate \ BR-SX-06 & QAOA/annealing: no advantage; SA dominates to $n = \mathbf{120}$ & negative \ BR-SX-08 & QML AUC **0.817** $\rightarrow$ **0.691**; sensing knee $\sigma \approx \mathbf{0.14}$ & negative \ longtable

tocsectionReferences

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 P I Ford, *KRONOS Physics De-Risking Register*, v8, Kronos Fusion Energy 2026 series. [doi:10.5281/zenodo.22133057](https://doi.org/10.5281/zenodo.22133057)
- 2 P I Ford, *Four Coupled Digital Twins for a Compact Spherical-Tokamak Breeder*, Kronos 2026 series. [doi:10.5281/zenodo.22645805](https://doi.org/10.5281/zenodo.22645805); <https://www.kronosfusionenergy.com/publications/four-coupled-digital-twins-for-a-compact-spherical-toka.html>
- 3 P I Ford, *A Certified Control-Barrier Safety Filter for Fusion Control*, Kronos 2026 series. [doi:10.5281/zenodo.22645716](https://doi.org/10.5281/zenodo.22645716); <https://www.kronosfusionenergy.com/publications/control-barrier-safety-clamp.html>
- 4 P I Ford, *Model-Predictive Burn Control Inside a Certified Safe Envelope*, Kronos 2026 series. [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>
- 5 P I Ford, *Neural-Operator Flux Surrogates for Control-Latency Transport*, Kronos 2026 series. [doi:10.5281/zenodo.22645803](https://doi.org/10.5281/zenodo.22645803); <https://www.kronosfusionenergy.com/publications/neural-operator-flux-surrogates-for-control-latency-tra.html>
- 6 P I Ford, *Sub-Millisecond Neural Equilibrium Reconstruction Robust to Sensor Dropout*, Kronos 2026 series. [doi:10.5281/zenodo.22645801](https://doi.org/10.5281/zenodo.22645801); <https://www.kronosfusionenergy.com/publications/sub-millisecond-neural-equilibrium-reconstruction-robus.html>
- 7 P I Ford, *Calibrated Uncertainty, Conformal Prediction, and Abstention for Safety-Critical Fusion Control*, Kronos 2026 series. [doi:10.5281/zenodo.22645813](https://doi.org/10.5281/zenodo.22645813); <https://www.kronosfusionenergy.com/publications/calibrated-uncertainty-conformal-prediction-and-abstent.html>
- 8 P I Ford, *Runtime Assurance with Signed-Ledger Provenance for Autonomous Fusion Control*, Kronos 2026 series. [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>
- 9 P I Ford, *Out-of-Distribution Detection and Epistemic Gating for Autonomous Fusion Control*, Kronos 2026 series. [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>
- 10 P I Ford, *Safe Reinforcement Learning from Operator Feedback for Fusion Control*, Kronos 2026 series. [doi:10.5281/zenodo.22645823](https://doi.org/10.5281/zenodo.22645823); <https://www.kronosfusionenergy.com/publications/safe-reinforcement-learning-from-operator-feedback-for-.html>
- 11 P I Ford, *Maturity-Gated Actuation Authority*, Kronos 2026 series. [doi:10.5281/zenodo.22645795](https://doi.org/10.5281/zenodo.22645795); <https://www.kronosfusionenergy.com/publications/maturity-gated-actuation-authority.html>
- 12 P I Ford, *The Plasma Diagnostic Suite and Real-Time State Estimation with a GNN Reconstructor*, Kronos 2026 series. [doi:10.5281/zenodo.22645943](https://doi.org/10.5281/zenodo.22645943); <https://www.kronosfusionenergy.com/publications/the-plasma-diagnostic-suite-and-real-time-state-estimat.html>
- 13 P I Ford, *A Real-Time Optimization Backbone for Fusion Plasma Control*, Kronos 2026 series. [doi:10.5281/zenodo.22645899](https://doi.org/10.5281/zenodo.22645899); <https://www.kronosfusionenergy.com/publications/a-real-time-optimization-backbone-for-fusion-plasma-con.html>
- 14 P I Ford, *MLOps for a Safety-Critical Fusion Control Stack*, Kronos 2026 series. [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>
- 15 P I Ford, *Global Sensitivity and Uncertainty Quantification across the Breeder and Burner Design Planes*, Kronos 2026 series. [doi:10.5281/zenodo.22645893](https://doi.org/10.5281/zenodo.22645893); <https://www.kronosfusionenergy.com/publications/global-sensitivity-and-uncertainty-quantification-acros.html>
- 16 P I Ford, *Bayesian Optimal Experimental Design and Multi-Fidelity Data Fusion*, Kronos 2026 series. [doi:10.5281/zenodo.22645895](https://doi.org/10.5281/zenodo.22645895); <https://www.kronosfusionenergy.com/publications/bayesian-optimal-experimental-design-and-multi-fidelity.html>

- 17 P I Ford, *Resource-Honest Quantum Computing for Fusion: Variational Landau-Damping and ADAPT-VQE Alloy Chemistry*, Kronos 2026 series. [doi:10.5281/zenodo.22645718](https://doi.org/10.5281/zenodo.22645718); <https://www.kronosfusionenergy.com/publications/quantum-computing-for-fusion.html>
- 18 P I Ford, *Quantum Optimization and Reinforcement Learning for Real-Time Plasma Control, without Claimed Advantage*, Kronos 2026 series. [doi:10.5281/zenodo.22645903](https://doi.org/10.5281/zenodo.22645903); <https://www.kronosfusionenergy.com/publications/quantum-optimization-and-reinforcement-learning-for-rea.html>
- 19 P I Ford, *Quantum Machine-Learning Surrogates, Entanglement-Enhanced Magnetometry and Fault-Tolerant Resource Estimates*, Kronos 2026 series. [doi:10.5281/zenodo.22645905](https://doi.org/10.5281/zenodo.22645905); <https://www.kronosfusionenergy.com/publications/quantum-machine-learning-surrogates-entanglement-enhanc.html>
- 20 P I Ford, *Verification-First Fusion Modeling: Frozen Anchors and a Byte-for-Byte Provenance Catalog*, Kronos 2026 series. [doi:10.5281/zenodo.22645710](https://doi.org/10.5281/zenodo.22645710); <https://www.kronosfusionenergy.com/publications/verification-and-validation-47-code-provenance.html>
- 21 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)
- 22 B N Sorbom *et al*, ARC: a compact, high-field, fusion nuclear science facility and demonstration power plant with demountable magnets, *Fusion Eng. Des.* **100** 378 (2015). [doi:10.1016/j.fusengdes.2015.07.008](https://doi.org/10.1016/j.fusengdes.2015.07.008)
- 23 A J Creely *et al*, Overview of the SPARC tokamak, *J. Plasma Phys.* **86** 865860502 (2020). [doi:10.1017/S0022377820001257](https://doi.org/10.1017/S0022377820001257)
- 24 M E Austin *et al*, Achievement of reactor-relevant performance in negative triangularity shape in the DIII-D tokamak, *Phys. Rev. Lett.* **122** 115001 (2019). [doi:10.1103/PhysRevLett.122.115001](https://doi.org/10.1103/PhysRevLett.122.115001)
- 25 A Marinoni, O Sauter and S Coda, A brief history of negative triangularity tokamak plasmas, *Rev. Mod. Plasma Phys.* **5** 6 (2021). [doi:10.1007/s41614-021-00054-0](https://doi.org/10.1007/s41614-021-00054-0)
- 26 F Troyon, R Gruber, H Saurenmann, S Semenzato and S Succi, MHD-limits to plasma confinement, *Plasma Phys. Control. Fusion* **26** 209 (1984). [doi:10.1088/0741-3335/26/1A/319](https://doi.org/10.1088/0741-3335/26/1A/319)
- 27 M Greenwald, Density limits in toroidal plasmas, *Plasma Phys. Control. Fusion* **44** R27 (2002). [doi:10.1088/0741-3335/44/8/201](https://doi.org/10.1088/0741-3335/44/8/201)
- 28 S Hahn *et al*, 45.5-tesla direct-current magnetic field generated with a high-temperature superconducting magnet, *Nature* **570** 496 (2019). [doi:10.1038/s41586-019-1293-1](https://doi.org/10.1038/s41586-019-1293-1)
- 29 L L Lao *et al*, 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)
- 30 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)
- 31 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)
- 32 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)
- 33 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)
- 34 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)
- 35 C Rea, K J Montes, K G Erickson, R S Granetz and R A Tinguely, 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)
- 36 D Q Mayne, J B Rawlings, C V Rao and P O M Sokaert, Constrained model predictive control: stability and optimality, *Automatica* **36** 789 (2000). [doi:10.1016/S0005-1098\(99\)00214-9](https://doi.org/10.1016/S0005-1098(99)00214-9)
- 37 D Q Mayne, Model predictive control: recent developments and future promise, *Automatica* **50** 2967 (2014). [doi:10.1016/j.automatica.2014.10.128](https://doi.org/10.1016/j.automatica.2014.10.128)
- 38 B Stellato, G Banjac, P Goulart, A Bemporad and S Boyd, OSQP: an operator splitting solver for quadratic programs, *Math. Program. Comput.* **12** 637 (2020). [doi:10.1007/s12532-020-00179-2](https://doi.org/10.1007/s12532-020-00179-2)

- 39 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** 3861 (2017). [doi:10.1109/TAC.2016.2638961](https://doi.org/10.1109/TAC.2016.2638961)
- 40 A D Ames *et al*, Control barrier functions: theory and applications, *European Control Conf. (ECC)* 3420 (2019). [doi:10.23919/ECC.2019.8796030](https://doi.org/10.23919/ECC.2019.8796030)
- 41 F Blanchini, Set invariance in control, *Automatica* **35** 1747 (1999). [doi:10.1016/S0005-1098\(99\)00113-2](https://doi.org/10.1016/S0005-1098(99)00113-2)
- 42 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)
- 43 G E Karniadakis *et al*, Physics-informed machine learning, *Nat. Rev. Phys.* **3** 422 (2021). [doi:10.1038/s42254-021-00314-5](https://doi.org/10.1038/s42254-021-00314-5)
- 44 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)
- 45 Z Li *et al*, Fourier neural operator for parametric partial differential equations, *ICLR* (2021). [doi:10.48550/arXiv.2010.08895](https://doi.org/10.48550/arXiv.2010.08895)
- 46 K L van de Plassche *et al*, Fast modeling of turbulent transport in fusion plasmas using neural networks, *Phys. Plasmas* **27** 022310 (2020). [doi:10.1063/1.5134126](https://doi.org/10.1063/1.5134126)
- 47 J Citrin *et al*, Real-time capable first principle based modelling of tokamak turbulent transport, *Nucl. Fusion* **55** 092001 (2015). [doi:10.1088/0029-5515/55/9/092001](https://doi.org/10.1088/0029-5515/55/9/092001)
- 48 O Meneghini *et al*, Integrated modeling applications for tokamak experiments with OMFIT, *Nucl. Fusion* **55** 083008 (2015). [doi:10.1088/0029-5515/55/8/083008](https://doi.org/10.1088/0029-5515/55/8/083008)
- 49 P Benner, S Gugercin and K Willcox, A survey of projection-based model reduction methods for parametric dynamical systems, *SIAM Rev.* **57** 483 (2015). [doi:10.1137/130932715](https://doi.org/10.1137/130932715)
- 50 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)
- 51 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)
- 52 J Candy and R E Waltz, An Eulerian gyrokinetic-Maxwell solver, *J. Comput. Phys.* **186** 545 (2003). [doi:10.1016/S0021-9991\(03\)00079-2](https://doi.org/10.1016/S0021-9991(03)00079-2)
- 53 F Jenko, W Dorland, M Kotschenreuther and B N Rogers, Electron temperature gradient driven turbulence, *Phys. Plasmas* **7** 1904 (2000). [doi:10.1063/1.874014](https://doi.org/10.1063/1.874014)
- 54 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)
- 55 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)
- 56 R E Kalman, A new approach to linear filtering and prediction problems, *J. Basic Eng.* **82** 35 (1960). [doi:10.1115/1.3662552](https://doi.org/10.1115/1.3662552)
- 57 G Evensen, The ensemble Kalman filter: theoretical formulation and practical implementation, *Ocean Dyn.* **53** 343 (2003). [doi:10.1007/s10236-003-0036-9](https://doi.org/10.1007/s10236-003-0036-9)
- 58 P L Houtekamer and H L Mitchell, Data assimilation using an ensemble Kalman filter technique, *Mon. Weather Rev.* **126** 796 (1998). [doi:10.1175/1520-0493\(1998\)126<0796:DAUAEK>2.0.CO;2](https://doi.org/10.1175/1520-0493(1998)126<0796:DAUAEK>2.0.CO;2)
- 59 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)
- 60 C L Degen, F Reinhard and P Cappellaro, Quantum sensing, *Rev. Mod. Phys.* **89** 035002 (2017). [doi:10.1103/RevModPhys.89.035002](https://doi.org/10.1103/RevModPhys.89.035002)
- 61 J F Barry *et al*, Sensitivity optimization for NV-diamond magnetometry, *Rev. Mod. Phys.* **92** 015004 (2020). [doi:10.1103/RevModPhys.92.015004](https://doi.org/10.1103/RevModPhys.92.015004)
- 62 L Rondin *et al*, Magnetometry with nitrogen-vacancy defects in diamond, *Rep. Prog. Phys.* **77** 056503 (2014). [doi:10.1088/0034-4885/77/5/056503](https://doi.org/10.1088/0034-4885/77/5/056503)
- 63 V Giovannetti, S Lloyd and L Maccone, Advances in quantum metrology, *Nat. Photonics* **5** 222 (2011). [doi:10.1038/nphoton.2011.35](https://doi.org/10.1038/nphoton.2011.35)

- 64 M Marchevsky, Quench detection and protection for high-temperature superconductor accelerator magnets, *Instruments* **5** 27 (2021). [doi:10.3390/instruments5030027](https://doi.org/10.3390/instruments5030027).
- 65 J Preskill, Quantum computing in the NISQ era and beyond, *Quantum* **2** 79 (2018). [doi:10.22331/q-2018-08-06-79](https://doi.org/10.22331/q-2018-08-06-79).
- 66 M Cerezo *et al*, Variational quantum algorithms, *Nat. Rev. Phys.* **3** 625 (2021). [doi:10.1038/s42254-021-00348-9](https://doi.org/10.1038/s42254-021-00348-9).
- 67 K Bharti *et al*, Noisy intermediate-scale quantum algorithms, *Rev. Mod. Phys.* **94** 015004 (2022). [doi:10.1103/RevModPhys.94.015004](https://doi.org/10.1103/RevModPhys.94.015004).
- 68 J R McClean, S Boixo, V N Smelyanskiy, R Babbush and H Neven, Barren plateaus in quantum neural network training landscapes, *Nat. Commun.* **9** 4812 (2018). [doi:10.1038/s41467-018-07090-4](https://doi.org/10.1038/s41467-018-07090-4).
- 69 A W Harrow, A Hassidim and S Lloyd, Quantum algorithm for linear systems of equations, *Phys. Rev. Lett.* **103** 150502 (2009). [doi:10.1103/PhysRevLett.103.150502](https://doi.org/10.1103/PhysRevLett.103.150502).
- 70 G H Low and I L Chuang, Hamiltonian simulation by qubitization, *Quantum* **3** 163 (2019). [doi:10.22331/q-2019-07-12-163](https://doi.org/10.22331/q-2019-07-12-163).
- 71 H R Grimsley, S E Economou, E Barnes and N J Mayhall, An adaptive variational algorithm for exact molecular simulations, *Nat. Commun.* **10** 3007 (2019). [doi:10.1038/s41467-019-10988-2](https://doi.org/10.1038/s41467-019-10988-2).
- 72 A Peruzzo *et al*, A variational eigenvalue solver on a photonic quantum processor, *Nat. Commun.* **5** 4213 (2014). [doi:10.1038/ncomms5213](https://doi.org/10.1038/ncomms5213).
- 73 A Kandala *et al*, Hardware-efficient variational quantum eigensolver for small molecules and quantum magnets, *Nature* **549** 242 (2017). [doi:10.1038/nature23879](https://doi.org/10.1038/nature23879).
- 74 Y Cao *et al*, Quantum chemistry in the age of quantum computing, *Chem. Rev.* **119** 10856 (2019). [doi:10.1021/acs.chemrev.8b00803](https://doi.org/10.1021/acs.chemrev.8b00803).
- 75 I M Georgescu, S Ashhab and F Nori, Quantum simulation, *Rev. Mod. Phys.* **86** 153 (2014). [doi:10.1103/RevModPhys.86.153](https://doi.org/10.1103/RevModPhys.86.153).
- 76 S Lee *et al*, Evaluating the evidence for exponential quantum advantage in ground-state quantum chemistry, *Nat. Commun.* **14** 1952 (2023). [doi:10.1038/s41467-023-37587-6](https://doi.org/10.1038/s41467-023-37587-6).
- 77 E Farhi, J Goldstone and S Gutmann, A quantum approximate optimization algorithm (2014). [doi:10.48550/arXiv.1411.4028](https://doi.org/10.48550/arXiv.1411.4028).
- 78 C Durr and P Hoyer, A quantum algorithm for finding the minimum (1996). [doi:10.48550/arXiv.quant-ph/9607014](https://doi.org/10.48550/arXiv.quant-ph/9607014).
- 79 A G Fowler, M Mariantoni, J M Martinis and A N Cleland, Surface codes: towards practical large-scale quantum computation, *Phys. Rev. A* **86** 032324 (2012). [doi:10.1103/PhysRevA.86.032324](https://doi.org/10.1103/PhysRevA.86.032324).
- 80 V Havlicek *et al*, Supervised learning with quantum-enhanced feature spaces, *Nature* **567** 209 (2019). [doi:10.1038/s41586-019-0980-2](https://doi.org/10.1038/s41586-019-0980-2).
- 81 A Engel, G Smith and S E Parker, Quantum algorithm for the Vlasov equation, *Phys. Rev. A* **100** 062315 (2019). [doi:10.1103/PhysRevA.100.062315](https://doi.org/10.1103/PhysRevA.100.062315).
- 82 I Y Dodin and E A Startsev, On applications of quantum computing to plasma simulations, *Phys. Plasmas* **28** 092101 (2021). [doi:10.1063/5.0056974](https://doi.org/10.1063/5.0056974).
- 83 I Joseph, Koopman–von Neumann approach to quantum simulation of nonlinear classical dynamics, *Phys. Rev. Research* **2** 043102 (2020). [doi:10.1103/PhysRevResearch.2.043102](https://doi.org/10.1103/PhysRevResearch.2.043102).

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–³He 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

AI and Quantum Computing for Fusion: ML Surrogates, Real-Time Control, and a Resource-Honest Quantum Frontier

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

[Zenodo · doi:10.5281/zenodo.22645702](https://doi.org/10.5281/zenodo.22645702)

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

[The Physics De-Risking Register](#)