

Physics-Informed Neural Networks for Real-Time State Estimation in High-Altitude Hydropower Cascades

Mohammad Salem¹, Mohammed Saif²

¹Department of Electrical Engineering, King Abdullah University of Science and Technology (KAUST), Thuwal, Saudi Arabia

²Department of Electrical and Computer Engineering, Prince Sultan University, Riyadh, Saudi Arabia

Article Info

Article history:

Received Dec 3, 2025

Revised Jan 17, 2026

Accepted Feb 7, 2026

Keywords:

physics-informed neural networks, hydraulic transients, water hammer, state estimation, hyperbolic systems, scientific machine learning, hydropower

ABSTRACT

Real-time state estimation of hydraulic transients in high-head hydropower cascades is a prerequisite for closed-loop governor control and for the fast frequency reserves that mountainous power systems increasingly rely on. Classical solvers for the water hammer equations, most notably the method of characteristics, are constrained by the Courant–Friedrichs–Lewy condition, which couples the temporal and spatial resolutions and forces the reconstruction of an entire trajectory even when only a handful of pointwise state values are required. This work develops a physics-informed neural network framework for the one-dimensional unsteady pipe flow system that governs penstock dynamics, and analyses it both theoretically and empirically. On the theoretical side we exploit the symmetric hyperbolic structure of the nondimensionalised system together with the monotonicity of the quadratic friction term to derive an a posteriori error estimate in which the growth of the bound is at most proportional to the square root of the horizon rather than exponential, and we combine this estimate with a quasi-Monte Carlo quadrature bound and a tanh approximation result to obtain a total error guarantee expressed entirely in terms of computable training residuals. The analysis identifies exact imposition of boundary data and characteristic-aligned domain decomposition as the two architectural choices that control the constants, and both are adopted in the method. On the empirical side we construct a reference solution for a 1200 m, 300 m gross head penstock using a method of characteristics solver operating at unit Courant number, verified to machine precision against the Joukowski and Michaud–Allievi limits and shown to be grid independent to eight significant figures. Against this reference the proposed model attains a relative L_2 error of 1.14% in piezometric head and 1.32% in discharge, degrades to 1.42% under 5% measurement noise and to 1.41% with two measurement channels, and recovers the acoustic wave speed to within 0.4% in the inverse setting. A parametric formulation amortises training across the closure-time family and supports online assimilation of supervisory telemetry in 6.4 ms. We also correct a claim that recurs in this literature: measured against a competently implemented reference solver the wall-clock advantage of the surrogate on full-field reconstruction is modest, and the operationally decisive advantages are instead constant-time pointwise evaluation, amortisation across operating scenarios, and end-to-end differentiability for control.

This is an open access article under the CC BY-SA license.



Corresponding Author:

Mohammad Salem

Department of Electrical Engineering

King Abdullah University of Science and Technology (KAUST)

Thuwal, Saudi Arabia

Email: mohammad.salem@kaust.edu.sa

NOMENCLATURE

Symbol	Definition	Units
$H(x, t)$	piezometric head	m
$Q(x, t)$	volumetric discharge	$\text{m}^3 \text{s}^{-1}$
a	acoustic wave speed in the fluid–pipe system	m s^{-1}
A	penstock cross-sectional area	m^2
D	internal penstock diameter	m
L	penstock length	m
f	Darcy–Weisbach friction factor	–
g	gravitational acceleration	m s^{-2}
T_c	guide-vane closure time	s
$\tau(t)$	dimensionless valve opening	–
h, q	nondimensional head and discharge perturbations	–
λ	nondimensional friction number	–
θ	network parameters	–
μ	operating-condition parameter vector	–
r_c, r_m	continuity and momentum residuals	–
$\mathcal{E}(t)$	error energy functional	–
$N_{col}, N_{ic}, N_{bc}, N_d$	numbers of collocation, initial, boundary and data points	–

1. INTRODUCTION

Hydropower cascades in high-altitude terrain operate with structural and hydraulic margins that are narrower than those of lowland installations. Gross heads of several hundred metres are developed over penstock lengths of one to two kilometres, so the static pressure is already close to the design envelope of the conduit before any transient occurs. Under rapid load rejection, guide-vane closure converts the momentum of the water column into a pressure wave whose amplitude, in the limiting case, is given by the Joukowsky relation $\Delta H = a\Delta V/g$ and can exceed the static head. The same event produces a downsurge on the reflected wave that may approach the vapour pressure and initiate column separation. The margin between acceptable and destructive operation is therefore set by the transient rather than the steady state, and the ability to estimate the transient state in real time is a direct determinant of how aggressively a plant may be dispatched.

This constraint has become economically material. Power systems in mountainous regions are absorbing increasing shares of variable renewable generation, and hydropower is the principal source of the fast frequency response that makes such absorption feasible. Exploiting that capability requires governor control loops that can act on a prediction of the pressure state rather than on a conservative envelope fixed at commissioning. Closed-loop model predictive control of guide vanes, waterway-aware unit commitment and online detection of anomalous conduit behaviour all require the same primitive: evaluation of head and discharge at arbitrary points of the penstock, at latencies compatible with supervisory control cycles of 10 to 100 ms.

The established numerical technology does not deliver this primitive in the form required. The method of characteristics transforms the hyperbolic system into ordinary differential equations along the characteristic curves $dx/dt = \pm a$ and integrates them on a staggered grid. Its stability and accuracy are governed by the Courant–Friedrichs–Lewy condition

$$\Delta t \leq \frac{\Delta x}{a},$$

and in the canonical implementation the condition is enforced as an equality so that the characteristics pass exactly through grid nodes. Numerical dispersion is then eliminated and the scheme is, in the frictionless limit, exact. This is a considerable strength, and any claim that the method of characteristics is inaccurate should be treated with suspicion. Its limitations are structural rather than numerical. First, the grid is rigid: changing the spatial resolution changes the time step and therefore the entire discrete trajectory, and interpolation is required whenever the state is needed away from a node. Second, the solver is sequential in time and reconstructs the full field even when a single point-wise value is wanted. Third, the solution is not differentiable with respect to the physical parameters or the control input in any form directly usable by a gradient-based controller, and obtaining such sensitivities requires an adjoint implementation of comparable complexity to the solver itself. Fourth, the solver assimilates no data: sparse and noisy supervisory measurements cannot be reconciled with the model except through an external filter whose linearisation is itself expensive for a distributed hyperbolic system.

Purely data-driven surrogates address the latency question and none of the others. Sequence models trained on simulated or historical transients interpolate acceptably inside their training distribution and violate mass and momentum conservation outside it, typically by damping wave amplitudes and accumulating phase error over successive reflections. Neither failure mode is detectable from the model output alone, which is disqualifying for a component in a protection-adjacent control loop.

Physics-informed neural networks occupy the intermediate position. The residual of the governing system, evaluated by automatic differentiation on the network output, is added to the training objective, so that the hypothesis space is restricted towards functions satisfying the conservation laws. The resulting estimator is mesh-free, differentiable in all its arguments by construction, and able to fuse sparse observations with the physical model in a single optimisation problem. A related physics-informed neural network framework, developed for pollutant transport in groundwater rather than for hydraulic transients, is given by Souleymane et al. [38]; the underlying residual-penalty idea is the same even though the application is quite different. The literature on such models for hyperbolic systems is, however, more equivocal than the literature on their elliptic and parabolic counterparts: convergence is slower, sharp fronts are approximated poorly, and reported failures are common enough that a claim of success requires both a theoretical account of when the method should work and an empirical protocol capable of detecting when it does not.

1.1. Contributions

This paper makes the following contributions.

1.2. Organisation

1. **A structure-exploiting error estimate.** We nondimensionalise the water hammer system into a symmetric hyperbolic form and show that the quadratic friction term is monotone. Combining the two yields an a posteriori bound in which the L_2 state error is controlled by the training residuals with a constant growing as $\sqrt{T}$. This is a qualitative improvement over the exponential constants that generic stability arguments produce for nonlinear systems, and it is what makes residual-based certification meaningful over the multi-period horizons of interest.
2. **A complete error decomposition.** The a posteriori estimate is combined with a Koksma–Hlawka quadrature bound for low-discrepancy collocation and with an approximation result for tanh networks, producing a guarantee in which every term is either computable after training or explicitly dependent on network width, depth and sample count.
3. **Theory-driven architecture.** The analysis shows that soft boundary penalties reintroduce an exponential constant through the boundary flux term, and that the achievable approximation rate is limited by the kinks that propagate along the characteristics emanating from the corner of the space-time domain. We therefore impose boundary and initial data exactly through a distance-function ansatz and decompose the domain along characteristic lines, and we quantify the effect of each choice.
4. **A parametric formulation with online assimilation.** Training is amortised over a parameter family covering closure time, initial discharge, wave speed and friction factor. At run time the operating point is identified from supervisory telemetry by a small Gauss–Newton problem in the parameter vector with the network weights frozen, which converts a surrogate that is expensive to train into an estimator that is cheap to condition on new data.
5. **A verified reference and an honest cost accounting.** The reference solver is verified to machine precision against the Joukowsky and Michaud–Allievi limits and shown to be grid independent, and its cost is measured rather than asserted. We find that the speedup of the surrogate over a competent reference implementation on full-field reconstruction is far smaller than is commonly reported, and we relocate the operational case for the surrogate to the properties that survive this correction.

Section 2 reviews the relevant literature. Section 3 derives the governing system, its characteristic structure and its nondimensional form. Section 4 specifies the estimator. Section 5 contains the theoretical analysis. Section 6 describes the experimental protocol and Section 7 the results. Section 8 discusses limitations and deployment considerations, and Section 9 concludes. Appendices contain the derivation of the governing equations, the reference solver, deferred proof details and the full hyperparameter specification.

2. RELATED WORK

Hydraulic transient modelling. The theory of pressure waves in closed conduits originates with the analyses of Joukowsky and Allievi and reached its modern computational form with the method of characteristics, developed for pipeline systems in the second half of the twentieth century and codified by Wylie and Streeter [1] and Chaudhry [2]. Ghidaoui et al. [3] survey the field and delineate the regimes in which the one-dimensional model is adequate. The principal refinement relevant here concerns friction. The quasi-steady Darcy–Weisbach term underestimates the attenuation of successive pressure peaks, and frequency-dependent corrections were introduced by Zielke [4]

for laminar flow and extended to turbulent conditions by Vardy and Brown [5]; instantaneous acceleration-based models due to Brunone and co-workers are reviewed by Bergant et al. [6]. A related Godunov-type scheme for moving steady states, developed for the Saint-Venant–Exner equations of sediment transport rather than for penstock transients, is given by Sore et al. [39] and illustrates a similar characteristic-respecting discretisation philosophy in an unrelated shallow-water setting. Finite volume and discontinuous Galerkin discretisations have been proposed where the rigidity of the characteristic grid is problematic, at the cost of introducing the numerical dispersion that the method of characteristics avoids.

Physics-informed learning. The residual-penalty formulation was introduced by Raissi et al. [7], with antecedents in the deep Galerkin method of Sirignano and Spiliopoulos [8], and is surveyed by Karniadakis et al. [9] and Cuomo et al. [10]. Practical training is dominated by the imbalance between loss terms; Wang et al. [11] characterise the resulting gradient pathologies and propose adaptive weighting, and Wang et al. [12] derive a weighting rule from the neural tangent kernel of the composite operator. Krishnapriyan et al. [13] document failure modes in convection-dominated regimes and show that curriculum and sequence-to-sequence strategies recover accuracy, a line continued by Wang et al. [14], who argue that the training dynamics must respect temporal causality and introduce the causal weighting scheme adopted here. Spectral bias towards low frequencies, characterised in the general setting by Rahaman et al. [15], is mitigated by the random Fourier feature embedding of Tancik et al. [16]. Domain decomposition variants due to Jagtap et al. [17] address both conditioning and the localisation of sharp features.

Theory. Shin et al. [18] establish consistency for linear elliptic and parabolic problems. Mishra and Molinaro [19] give generalisation error estimates in which the total error is bounded by the training error plus a quadrature term, the framework we follow. Approximation rates for tanh networks in Sobolev norms, which supply the remaining term, are due to De Ryck et al. [20]. For hyperbolic problems the picture is less complete; De Ryck and Mishra [21] treat several first-order systems, and the general obstruction, that stability constants for nonlinear hyperbolic systems degrade exponentially in time, is standard [22, 23]. The contribution of Section 5 is to show that for the water hammer system this obstruction is absent.

Operator learning and power system applications. DeepONet [24] and the Fourier neural operator [25] learn solution operators rather than individual solutions and are natural baselines for the parametric formulation of Section 4.6. Within power systems, Misyris et al. [26] applied physics-informed networks to swing dynamics, and Stiasny et al. [27] to transient stability assessment. In fluid conduits, Kissas et al. [28] applied the approach to one-dimensional blood flow, a system with the same hyperbolic structure as the present one, which is the closest methodological precedent to this work. Applications to leak localisation in water distribution networks have also appeared. To our knowledge the combination of a hyperbolic-structure-specific error estimate, a parametric formulation with online assimilation, and a verified cost comparison has not previously been presented for penstock transients.

3. PROBLEM FORMULATION

3.1. Governing System

Consider a penstock of length L , constant internal diameter D and cross-sectional area $A = \pi D^2/4$, conveying water from an upstream reservoir at $x = 0$ to a turbine inlet valve at $x = L$. The derivation in Appendix A applies conservation of mass and momentum to a control volume of a slightly compressible fluid in a linearly elastic conduit under the following assumptions.

A1. The flow is one-dimensional, with velocity and pressure uniform over each cross-section. **A2.** The conduit is prismatic, its wall is linearly elastic, and the combined fluid and wall elasticity is summarised by a constant wave speed a . **A3.** Convective acceleration is negligible relative to local acceleration. This is justified when the Mach number V/a is small; for the configuration of Section 6 it is 2.5×10^{-3} . **A4.** Head losses follow a quasi-steady Darcy–Weisbach law with constant friction factor f . Section 8 discusses the consequences of relaxing this assumption. **A5.** The absolute pressure remains above the vapour pressure throughout, so the flow is single-phase and column separation does not occur. This is verified a posteriori in Section 6.3.

Under A1–A5 the state (H, Q) satisfies, on $\Omega = (0, L) \times (0, T)$,

$$\begin{aligned} \frac{\partial H}{\partial t} + \frac{a^2}{gA} \frac{\partial Q}{\partial x} &= 0, \\ \frac{\partial Q}{\partial t} + gA \frac{\partial H}{\partial x} + \frac{f}{2DA} Q|Q| &= 0. \end{aligned}$$

The friction term is written with $Q|Q|$ rather than Q^2 so that the resistance opposes the flow under reverse motion, which occurs during the downsurge phase. This choice is not cosmetic: it makes the term monotone, and Section 5 shows that monotonicity is precisely what removes the exponential constant from the error estimate. A related structural argument exploiting monotone multivalued friction terms, developed for a three-dimensional Navier–Stokes system rather than for one-dimensional pipe flow, is given by Zongo and Kyelem [36]; the analogy is at the level of proof technique only.

The system is closed by the initial condition given by the steady state at discharge Q_0 ,

$$Q(x, 0) = Q_0, \quad H(x, 0) = H_r - \frac{fV_0^2}{2gD}x, \quad V_0 = Q_0/A,$$

by the upstream condition $H(0, t) = H_r$ representing a reservoir of effectively infinite storage, and by the downstream valve characteristic

$$Q(L, t) = \tau(t) Q_0 \sqrt{\frac{H(L, t)}{H_0^v}}, \quad H_0^v = H(L, 0),$$

where τ decreases from unity to zero over the closure interval $[0, T_c]$. The square-root dependence on the instantaneous head is the orifice characteristic of the guide-vane assembly and makes the downstream condition nonlinear; it is retained because linearising it changes the peak pressure by several per cent, as quantified in Section 6.3.

3.2. Characteristic Structure and Well-Posedness

Writing $u = (H, Q)^\top$ the system takes the form $u_t + Au_x + s(u) = 0$ with

$$A = \begin{pmatrix} 0 & a^2/(gA) \\ gA & 0 \end{pmatrix}, \quad s(u) = \left(0, \frac{f}{2DA}Q|Q|\right)^\top.$$

The eigenvalues of A are $\pm a$ with left eigenvectors $(gA, \pm a)$, so the system is strictly hyperbolic and the Riemann invariants are $gAH \pm aQ$. Since one eigenvalue is positive and one negative, exactly one scalar condition is required at each end of the interval, which is what the reservoir and valve conditions supply, and the initial-boundary value problem is well posed in the sense of Kreiss for C^1 data. Existence and uniqueness arguments of this kind ultimately rest on fixed-point machinery; a broader review of that machinery, in application areas unrelated to hyperbolic conservation laws, is given by Kamalov and Leung [41]. Because A is constant the characteristics are the straight lines $x \mp at = \text{const}$, and the regularity of the solution is transported along them without smoothing. This last observation is central to Section 5.4: any lack of smoothness in the boundary data, including the discontinuity in $\dot{\tau}$ at $t = 0$ and $t = T_c$, propagates undamped into the interior and reflects, so the solution is only piecewise smooth on a lattice of characteristic lines regardless of how smooth the data are elsewhere. No parabolic regularisation is available to hide this, and it caps the approximation rate attainable by a globally smooth network.

3.3. Nondimensionalisation

The dimensional system is badly scaled for gradient-based training: H is of order 10^2 , Q of order 10, and the coefficients $a^2/(gA)$ and gA differ by five orders of magnitude. We therefore rescale. Let

$$\tilde{x} = \frac{x}{L}, \quad \tilde{t} = \frac{at}{L}, \quad h = \frac{H - H_r}{\Delta H_J}, \quad q = \frac{Q}{Q_0}, \quad \Delta H_J = \frac{aQ_0}{gA},$$

where ΔH_J is the Joukowski head rise. Substitution gives, after dropping tildes,

$$h_t + q_x = 0, \quad q_t + h_x + \lambda q|q| = 0, \quad \lambda = \frac{fLV_0}{2Da},$$

on the unit square in space and $[0, \tilde{T}]$ in time with $\tilde{T} = aT/L$. The rescaled system is symmetric hyperbolic with unit characteristic speeds, both state components are of order unity, and the single remaining parameter λ measures friction against inertia. For the configuration of Section 6, $\lambda = 5.625 \times 10^{-3}$, confirming that the transient is weakly damped and that attenuation is a secondary effect over the horizon of interest. All network inputs and outputs are in these variables; all reported errors are relative and therefore invariant to the rescaling.

The nondimensional friction term inherits the property that matters. Writing $\phi(s) = s|s|$, for any $s, \sigma \in \mathbb{R}$,

$$(\phi(s) - \phi(\sigma))(s - \sigma) \geq 0,$$

with equality only when $s = \sigma$. The map is monotone but not globally Lipschitz, and it is the monotonicity, not any Lipschitz constant, that the analysis uses.

3.4. The Estimation Problem

Let $\mu \in \mathcal{M} \subset \mathbb{R}^p$ collect the operating parameters that vary between events, namely closure time, initial discharge, wave speed and friction factor. Instrumentation consists of pressure transducers at a small number of locations $\{x_k\}$ and a discharge measurement at the turbine inlet, sampled at the supervisory rate and corrupted by noise. The task is to construct a map

$$\mathcal{S} : (x, t, \mu) \longmapsto (h, q)$$

that can be evaluated at arbitrary (x, t) in bounded time, that is consistent with the governing system, that can be conditioned on the available measurements, and that is differentiable with respect to t and μ so that a controller may use its sensitivities. The remainder of the paper constructs and analyses such a map.

4. METHOD

4.1. Constrained Ansatz

Let $N_h, N_q : \mathbb{R}^2 \rightarrow \mathbb{R}$ denote the two outputs of a multilayer perceptron with parameters θ . Rather than represent h and q directly by the network, we compose the network with functions that satisfy the initial condition and the upstream boundary condition identically:

$$h_\theta(x, t) = -\lambda x + x \beta(t) N_h(x, t; \theta), \quad q_\theta(x, t) = 1 + \beta(t) N_q(x, t; \theta), \quad \beta(t) = 1 - e^{-t}.$$

The factor x enforces $h_\theta(0, t) = 0$ for every θ , which is the reservoir condition in nondimensional variables. The factor β vanishes at $t = 0$ with unit derivative and saturates thereafter, so $h_\theta(x, 0) = -\lambda x$ and $q_\theta(x, 0) = 1$ reproduce the steady state exactly while leaving the network unconstrained away from the initial line. Two of the four constraint sets are thereby removed from the objective. A related architectural technique for embedding constraints directly into a network's output, developed for fractional differential equations rather than for hyperbolic conservation laws, is the adaptive hybrid B-spline PINN of Padhy et al. [35]; the constraint-embedding principle is shared even though the equation class and the specific construction differ. The valve characteristic at $x = 1$ couples h and q nonlinearly and cannot be eliminated by an equally simple ansatz; it is retained as a penalty, and Section 5.3 quantifies the price of doing so.

4.2. Input Embedding

The transient contains a fundamental component at the wave return period $4L/a$, which is 4 in nondimensional time, together with harmonics generated by the reflection structure and by the closure law. A plain coordinate input biases the network towards the lowest of these components. We therefore precede the network with a random Fourier embedding,

$$\gamma(x, t) = [\sin(2\pi Bz), \cos(2\pi Bz)], \quad z = (x, t)^\top, \quad B_{ij} \sim \mathcal{N}(0, \sigma_F^2),$$

with m frequency pairs, so that the effective input dimension is $2m$. The scale σ_F is not tuned by grid search but read off the data: the closure law introduces energy up to a nondimensional frequency of order $1/\bar{T}_c$, and the reflection structure up to roughly the tenth harmonic of the fundamental, giving $\sigma_F \approx 3$. Section 7.6 confirms that performance is flat for $\sigma_F \in [2, 5]$ and degrades outside this range, the failure at large σ_F taking the characteristic form of a high-frequency ripple superposed on an otherwise correct wave.

4.3. Residuals and Objective

The residual operators are evaluated on the constrained ansatz by automatic differentiation,

$$r_c[\theta](x, t) = \partial_t h_\theta + \partial_x q_\theta, \quad r_m[\theta](x, t) = \partial_t q_\theta + \partial_x h_\theta + \lambda q_\theta |q_\theta|,$$

and the valve residual is

$$r_v[\theta](t) = q_\theta(1, t) - \tau(t) \sqrt{\max\{h_\theta(1, t)/h_0^v + 1, 0\}},$$

written in the nondimensional head with the clamp guarding against the square root of a negative argument during early training, when the network may briefly predict heads below the vapour limit. The objective is

$$\mathcal{L}(\theta) = w_r \mathcal{L}_r(\theta) + w_v \mathcal{L}_v(\theta) + w_d \mathcal{L}_d(\theta),$$

with

$$\mathcal{L}_r = \frac{1}{N_{col}} \sum_{j=1}^{N_{col}} c_j \left(|r_c(x_j, t_j)|^2 + |r_m(x_j, t_j)|^2 \right), \quad \mathcal{L}_v = \frac{1}{N_{bc}} \sum_{k=1}^{N_{bc}} |r_v(t_k)|^2,$$

$$\mathcal{L}_d = \frac{1}{N_d} \sum_{i=1}^{N_d} \left(|h_\theta(x_i, t_i) - h_i^*|^2 + \rho |q_\theta(x_i, t_i) - q_i^*|^2 \right),$$

where ρ balances the two measured quantities and is set to unity in the nondimensional variables. The initial-condition and upstream terms present in the conventional formulation are absent by construction. The weights c_j implement causal reweighting and are defined in Section 4.5.

4.4. Collocation

Collocation points are drawn from a scrambled Sobol sequence on $[0, 1] \times [0, \tilde{T}]$ rather than by Latin hypercube sampling. The reason is analytical rather than empirical: the quadrature term in the error bound of Section 5.2 is controlled by the star discrepancy of the point set through the Koksma–Hlawka inequality, which for a Sobol sequence in dimension d decays as $(\log N)^d/N$, against $N^{-1/2}$ for a Latin hypercube design in the absence of additional structure. In two and three dimensions the difference is substantial at the sample counts used here. The point set is resampled every 10^3 iterations to avoid overfitting to a fixed design.

4.5. Loss Balancing and Temporal Causality

Two pathologies dominate the optimisation. The first is scale imbalance between residual and data terms, whose gradients differ by orders of magnitude at initialisation. We adopt the neural tangent kernel rule: writing $K_k(\theta)$ for the kernel of loss term k , the weights are updated every M iterations as

$$w_k \leftarrow (1 - \alpha)w_k + \alpha \frac{\sum_l \text{Tr } K_l(\theta)}{\text{Tr } K_k(\theta)},$$

with $\alpha = 0.1$ and $M = 10^3$, and the traces estimated from a random subset of collocation points. This equalises the convergence rates of the terms rather than their magnitudes.

The second pathology is specific to evolution problems. Gradient descent on a residual aggregated uniformly over the space-time domain reduces the residual at late times before the solution at early times has converged, which allows the network to satisfy the equations along a trajectory that does not originate at the correct initial state. We therefore apply causal weights on a partition of $[0, \tilde{T}]$ into N_s bins,

$$c_j = \exp \left(-\epsilon \sum_{i < \iota(j)} \mathcal{L}_r^{(i)} \right),$$

where $\iota(j)$ is the bin containing t_j and $\mathcal{L}_r^{(i)}$ is the mean residual in bin i . A bin is effectively activated only once the residual in all preceding bins has fallen, which recovers the sequential structure of the underlying evolution without abandoning the global optimisation. We use $N_s = 32$ and increase ϵ through $\{10^{-2}, 10^{-1}, 1, 10\}$ on a schedule, reporting for each run the largest ϵ at which all bin weights exceed 0.99, which serves as a convergence certificate.

4.6. Characteristic-Aligned Decomposition

Section 3.2 established that the solution is piecewise smooth on a lattice of characteristic lines rather than globally smooth. A single network with a smooth activation must therefore either smear the kinks or expend capacity resolving them, and in the fast-closure regime this is the dominant error source. We partition the space-time domain along the characteristics emanating from the corners $(1, 0)$ and $(1, \tilde{T}_c)$ and their reflections at $x = 0$ and $x = 1$, and assign an independent subnetwork to each cell. Continuity of h and q across an interface Γ is imposed by the penalty

$$\mathcal{L}_\Gamma = \frac{1}{N_\Gamma} \sum_\Gamma \left(\|h_\theta^+ - h_\theta^- \|^2 + \|q_\theta^+ - q_\theta^- \|^2 \right),$$

which enforces C^0 matching only, since the exact solution is generally not C^1 across these lines. For the primary case of Section 6 the decomposition is inactive by default, since the closure is slow enough that the fronts are mild; it is enabled for the fast-closure study of Section 7.7, where it reduces the error by a factor of three.

For long horizons we additionally partition time into slabs of nondimensional length 2, matching the wave return period $2L/a$, and train them sequentially with the terminal state of each slab supplying the initial condition of the next through the ansatz of Section 4.1. This bounds the effective horizon entering the error estimate and is the practical realisation of the $\sqrt{T}$ scaling derived in Section 5.1.

4.7. Parametric Formulation

A network trained for a single closure event is a solution, not an estimator. We therefore append the operating parameters to the input, writing

$$\hat{\mu} = \left(\frac{T_c - \tilde{T}_c}{\sigma_{T_c}}, \frac{Q_0 - \tilde{Q}_0}{\sigma_{Q_0}}, \frac{a - \bar{a}}{\sigma_a}, \frac{f - \bar{f}}{\sigma_f} \right)$$

and $N_h, N_q : (x, t, \hat{\mu}) \mapsto \mathbb{R}$. The residual, valve and data losses are integrated over $\mathcal{M}$ as well as over the space-time domain, with the joint point set again drawn from a scrambled Sobol sequence in $2 + p$ dimensions. Note that λ and $\tilde{T}_c$ are functions of μ , so the nondimensional system itself varies over the family; this is handled by evaluating the coefficients pointwise from the sampled μ rather than fixing them.

4.8. Online Assimilation

At run time the operating point is unknown: the guide-vane schedule is commanded and therefore known, but the effective wave speed drifts with entrained air content and wall condition, and the friction factor drifts with surface roughness. Retraining is out of the question at supervisory latencies. Instead we freeze θ and solve for μ alone,

$$\hat{\mu}(t) = \arg \min_{\mu \in \mathcal{M}} \sum_{i \in \mathcal{W}(t)} \|y_i - \mathcal{H}[h_\theta, q_\theta](x_i, t_i, \mu)\|_{R^{-1}}^2 + \|\mu - \mu_{prior}\|_{P^{-1}}^2,$$

over a trailing window $\mathcal{W}(t)$ of telemetry, where $\mathcal{H}$ is the observation operator and R, P are the measurement and prior covariances. Because p is small and the network is differentiable with respect to μ , this is a p -dimensional nonlinear least-squares problem solved by Gauss–Newton with Jacobians obtained by automatic differentiation. Convergence takes a handful of iterations from the previous estimate as a warm start. The full estimator is then $(x, t) \mapsto (h_\theta, q_\theta)(x, t, \hat{\mu})$, evaluable at arbitrary points at the cost of a single forward pass.

4.9. Inverse Identification

The same machinery solves the offline identification problem. Treating a and f as trainable scalars alongside θ and minimising the same objective on measured data recovers the physical parameters as a by-product of fitting the state. No modification of the residual is required, which is the standard argument for the approach and one of the few places where it is clearly superior to a classical solver.

4.10. Training Procedure

Algorithm 1. *Offline training.*

Input: parameter domain $\mathcal{M}$, horizon $\tilde{T}$, budgets N_{col}, N_{bc}, N_d .

1. Draw a scrambled Sobol point set on $[0, 1] \times [0, \tilde{T}] \times \mathcal{M}$.
2. Initialise θ by the Glorot scheme; sample B and hold it fixed.
3. **Phase 1.** Minimise $\mathcal{L}$ by Adam for 2×10^4 iterations, learning rate 10^{-3} with cosine decay to 10^{-5} , updating w_k every 10^3 iterations by the rule of Section 4.5 and advancing the causal schedule when the certificate is met. Resample collocation points every 10^3 iterations.
4. **Phase 2.** Minimise $\mathcal{L}$ by L-BFGS in double precision with a strong Wolfe line search, fixed weights, fixed point set, terminating on a gradient norm below 10^{-8} or 1.5×10^4 iterations.
5. **Output:** θ^* , and the residual norms required by the certificate of Section 5.5.

The two-phase schedule is standard and its rationale is that Adam navigates the loss landscape efficiently but stalls at the accuracy required for stiff residuals, whereas L-BFGS exploits the curvature of the nearly quadratic basin reached at the end of Phase 1 but diverges if started far from it. The switch is made on a plateau criterion rather than at a fixed iteration count. Hyperparameter and architecture choices here are fixed rather than selected by a formal model-selection criterion; a related regularization-based approach to model selection, developed outside the physics-informed setting, is given by Kamalov et al. [43] and could in principle be adapted to choosing among candidate architectures.

5. THEORETICAL ANALYSIS

This section bounds the error of the trained estimator by quantities computable after training. Throughout, (h, q) denotes the classical solution of the nondimensional system on $\Omega_{\tilde{T}} = [0, 1] \times [0, \tilde{T}]$, assumed to lie in $C^1(\Omega_{\tilde{T}})$ and to satisfy $\sup |q| \leq \bar{q}$ and $h/h_0^v + 1 \geq \eta > 0$; the latter is the no-cavitation condition A5 and is verified numerically in Section 6.3. The approximation (h_θ, q_θ) is the ansatz of Section 4.1, which satisfies the initial condition and the upstream boundary condition identically and is C^1 because the activation is. We write $e = h_\theta - h$ and $\varepsilon = q_\theta - q$ for the pointwise errors, r_c, r_m for the interior residuals and r_v for the valve residual, and set

$$\mathcal{E}(t) = \frac{1}{2} \int_0^1 (e^2 + \varepsilon^2) dx, \quad R(t) = \|r_c(\cdot, t)\|_{L^2(0,1)}^2 + \|r_m(\cdot, t)\|_{L^2(0,1)}^2.$$

Finally $M = \sup_{t \leq \tilde{T}} |e(1, t)|$, which is finite because both the network output and the exact solution are bounded.

5.1. A Posteriori Stability

Theorem 1. For every $t \in [0, \tilde{T}]$,

$$\|e(\cdot, t)\|_{L^2}^2 + \|\varepsilon(\cdot, t)\|_{L^2}^2 \leq 2t(\|r_c\|_{L^2(\Omega_t)}^2 + \|r_m\|_{L^2(\Omega_t)}^2) + 8M\|r_v\|_{L^1(0,t)}.$$

Proof. Subtracting the governing equations from the residual definitions gives

$$e_t + \varepsilon_x = r_c, \quad \varepsilon_t + e_x + \lambda(\phi(q_\theta) - \phi(q)) = r_m,$$

with $\phi(s) = s|s|$. Multiplying the first by e , the second by ε , adding and integrating over $(0, 1)$,

$$\mathcal{E}'(t) = \int_0^1 (e r_c + \varepsilon r_m) dx - \int_0^1 \partial_x(e\varepsilon) dx - \lambda \int_0^1 \varepsilon(\phi(q_\theta) - \phi(q)) dx.$$

The last integrand is nonnegative pointwise, since $\varepsilon = q_\theta - q$ and ϕ is monotone, so the third term is bounded above by zero and may be discarded. This is the only place where the structure of the friction law is used, and it is used without any Lipschitz constant; a term modelled as λq^2 rather than $\lambda q|q|$ would fail to be monotone under flow reversal and the argument would not close.

For the second term, $\int_0^1 \partial_x(e\varepsilon) dx = e(1, t)\varepsilon(1, t) - e(0, t)\varepsilon(0, t)$, and the upstream contribution vanishes identically because the ansatz enforces $h_\theta(0, t) = 0$. Hence

$$\mathcal{E}'(t) \leq \int_0^1 (e r_c + \varepsilon r_m) dx - e(1, t)\varepsilon(1, t).$$

The interior term is estimated by the Cauchy–Schwarz inequality in the product space,

$$\int_0^1 (e r_c + \varepsilon r_m) dx \leq (\|e\|^2 + \|\varepsilon\|^2)^{1/2} R(t)^{1/2} = \sqrt{2\mathcal{E}(t)} R(t)^{1/2}.$$

For the boundary term write $b(t) = |r_v(t)|$ and decompose $\varepsilon(1, t) = r_v(t) + \tau(t)[\psi(h_\theta(1, t)) - \psi(h(1, t))]$ with $\psi(\cdot) = \sqrt{\cdot/h_0^v + 1}$. Since ψ is increasing and $h_\theta(1, t) - h(1, t) = e(1, t)$, the bracket carries the sign of $e(1, t)$, so the corresponding contribution to $-e(1, t)\varepsilon(1, t)$ is nonpositive and may be discarded. The remaining contribution is $-e(1, t)r_v(t) \leq Mb(t)$. Therefore

$$\mathcal{E}'(t) \leq \sqrt{2\mathcal{E}(t)} R(t)^{1/2} + Mb(t).$$

Fix $\delta > 0$ and set $y_\delta = (\mathcal{E} + \delta^2)^{1/2}$. Then $y'_\delta = \mathcal{E}'/(2y_\delta)$ and, using $\sqrt{2\mathcal{E}} \leq \sqrt{2}y_\delta$ and $1/y_\delta \leq 1/\delta$,

$$y'_\delta(t) \leq \frac{1}{\sqrt{2}} R(t)^{1/2} + \frac{Mb(t)}{2\delta}.$$

Integrating from 0 to t and using $\mathcal{E}(0) = 0$, which holds because the ansatz reproduces the initial state exactly, gives

$$\mathcal{E}(t)^{1/2} \leq y_\delta(t) \leq \delta + \frac{1}{\sqrt{2}} \int_0^t R^{1/2} ds + \frac{M}{2\delta} \int_0^t b ds.$$

The right-hand side is minimised at $\delta^* = (\frac{M}{2} \int_0^t b)^{1/2}$, whence

$$\mathcal{E}(t)^{1/2} \leq \frac{1}{\sqrt{2}} \int_0^t R^{1/2} ds + (2M \|r_v\|_{L^1(0,t)})^{1/2}.$$

Squaring, applying $(u + v)^2 \leq 2u^2 + 2v^2$, using $(\int_0^t R^{1/2})^2 \leq t \int_0^t R$ by the Cauchy–Schwarz inequality and $\|e\|^2 + \|\varepsilon\|^2 = 2\mathcal{E}$ yields the claim. ■

Three features of this estimate deserve comment.

The dependence on the horizon is $\sqrt{t}$, not exponential. For a general nonlinear hyperbolic system the standard argument absorbs the interior term with Young's inequality, producing $\mathcal{E}' \leq \mathcal{E} + \frac{1}{2}R$ and a Grönwall factor e^t . Here the symmetry of the flux matrix makes the interior flux an exact divergence and the monotonicity of the friction term makes the source dissipative, so no term proportional to $\mathcal{E}$ is generated and the estimate closes at the level of $\sqrt{\mathcal{E}}$. Over the horizon of interest, $\tilde{T} = 8.33$, the difference between $\sqrt{2t} \approx 4.1$ and $e^t \approx 4 \times 10^3$ is the difference between a bound that certifies and a bound that does not. See also Obogi and Evans for a related spectral-gap argument concerning approximate invariance in Hilbert spaces, developed outside the present PDE setting, which shares the general flavour of bounding a deviation quantity by a structurally controlled constant [37].

The two residual types enter with different strengths. Interior residuals appear through their L^2 norm and boundary residuals through the square root of their L^1 norm, so a boundary residual of size ϵ contributes $O(\sqrt{\epsilon})$ to the error while an interior residual of the same size contributes $O(\epsilon)$. Boundary conditions that can be imposed exactly should therefore be imposed exactly. This is the analytical justification for the ansatz of Section 4.1 and it is a quantitative statement, not a preference.

Had the upstream condition been imposed by penalty instead, the term $e(0, t)\varepsilon(0, t)$ would not vanish and would have to be controlled through a trace inequality, requiring H^1 control of the error and reintroducing a Grönwall factor. The following makes this precise.

Proposition 2. If the upstream condition is imposed by penalty with residual b_0 rather than exactly, then $\mathcal{E}(t) \leq e^t (\frac{1}{2} \int_0^t R + M \|b_0\|_{L^1(0,t)} + M \|r_v\|_{L^1(0,t)})$, and the exponential factor cannot in general be removed.

Proof. The boundary term now contributes $|e(0, t)\varepsilon(0, t)| \leq M|b_0(t)|$ with no sign, and the interior term is absorbed by Young's inequality as $\sqrt{2\mathcal{E}}R^{1/2} \leq \mathcal{E} + \frac{1}{2}R$, giving $\mathcal{E}' \leq \mathcal{E} + \frac{1}{2}R + M(|b_0| + b)$. Grönwall's inequality gives the stated bound. The factor is not removable because the argument no longer controls $\mathcal{E}'$ by a quantity vanishing with $\mathcal{E}$: the upstream boundary is an inflow boundary for the characteristic family travelling into the domain, and an error there is transported rather than dissipated. ■

5.2. From Continuous to Empirical Residuals

Theorem 1 involves the L^2 norms of the residuals over the domain, whereas training minimises Monte Carlo estimates of those norms. The gap is a quadrature error, and because the residual of a tanh network is smooth it is controlled by discrepancy rather than by variance. Let $\mathcal{Z} = \{z_j\}_{j=1}^{N_{col}} \subset [0, 1]^d$ with $d = 2$ in the single-scenario case and $d = 2 + p$ in the parametric case, and let $D_N^*(\mathcal{Z})$ denote its star discrepancy. Write $F_\theta = |r_c|^2 + |r_m|^2$ and let $V(F_\theta)$ be its variation in the sense of Hardy and Krause.

Theorem 3. With $|\Omega|$ the measure of the rescaled domain,

$$\|r_c\|_{L^2(\Omega)}^2 + \|r_m\|_{L^2(\Omega)}^2 \leq |\Omega| \left(\frac{1}{N_{col}} \sum_{j=1}^{N_{col}} F_\theta(z_j) + V(F_\theta) D_{N_{col}}^*(\mathcal{Z}) \right),$$

and if $\mathcal{Z}$ is the first N_{col} points of a Sobol sequence then $D_{N_{col}}^* = O(N_{col}^{-1}(\log N_{col})^{d-1})$.

Proof. After affine rescaling of Ω to the unit cube the first statement is the Koksma–Hlawka inequality applied to F_θ , and the second is the standard discrepancy bound for digital (t, d) -sequences. ■

The bound is vacuous unless $V(F_\theta)$ is finite and controlled. It is, and the control is explicit in the network parameters.

Proposition 4. Let N_θ be a tanh network of depth L with weight matrices $W_1, \dots, W_L$. Then all mixed partial derivatives of F_θ of order at most d are bounded on Ω by a constant depending only on $d, L, \lambda, \bar{q}$ and $\prod_l \|W_l\|_\infty$, and consequently $V(F_\theta) < \infty$.

Proof sketch. The Hardy–Krause variation is bounded by the sum of the L^1 norms of the mixed partials of order at most d , so it suffices to bound those. Proceed by induction on depth. The derivatives of tanh of every order are bounded by unity in absolute value on $\mathbb{R}$, and the pre-activation of layer l is an affine function of the output of layer $l - 1$. By the Faà di Bruno formula the mixed partials of order k of the layer- l output are finite sums of products of derivatives of tanh with mixed partials of the pre-activation of order at most k , each carrying one factor of $\|W_l\|_\infty$ per differentiation. Induction gives a bound of the form $C(k, L) \prod_l \|W_l\|_\infty^k$. The residuals are first-order differential expressions in the network plus the term $\lambda\phi(q_\theta)$, and ϕ is C^1 with $|\phi'| \leq 2\bar{q}$ on the relevant range, so the same bound applies to r_c, r_m with k increased by one, and to F_θ by the product rule. ■

Since Phase 2 of Algorithm 1 does not regularise the weights explicitly, $\prod_l \|W_l\|_\infty$ is monitored during training and reported, so that the constant in Theorem 3 is an observed rather than an assumed quantity.

5.3. Approximation

The remaining question is whether a network of the stated architecture can drive the residuals to zero at all. The following is not new and is recalled for completeness in the form given by De Ryck, Lanthaler and Mishra. A related approximation-theoretic result for a different operator family, namely new Baskakov operators, is given by Fadhil and Hassan [40]; it is cited here only as an indication that quantitative approximation results of this flavour recur across otherwise unrelated operator and network classes.

Proposition 5. For $u \in H^k([0, 1]^d)$ with $k \geq 3$ and every $N \in \mathbb{N}$ there exists a tanh network with two hidden layers of widths $O(N^d)$ such that $\|u - u_\theta\|_{H^1} \leq C N^{-(k-1)}$, with C depending on k, d and $\|u\|_{H^k}$.

Corollary 6. If in addition $(h, q) \in H^k(\Omega)$ with $k \geq 3$, there is a network of the above architecture for which $\|r_c\|_{L^2} + \|r_m\|_{L^2} \leq C' N^{-(k-1)}$, with C' depending additionally on λ and $\bar{q}$.

Proof. The interior residuals are first-order differential operators applied to the error plus the friction discrepancy. Hence $\|r_c\|_{L^2} \leq \|h_\theta - h\|_{H^1} + \|q_\theta - q\|_{H^1}$ and $\|r_m\|_{L^2} \leq \|q_\theta - q\|_{H^1} + \|h_\theta - h\|_{H^1} + \lambda \|\phi(q_\theta) - \phi(q)\|_{L^2}$. On the bounded set where $|q|, |q_\theta| \leq \bar{q}$ the map ϕ is Lipschitz with constant $2\bar{q}$, so the last term is bounded by $2\lambda\bar{q}\|q_\theta - q\|_{L^2}$. Applying Proposition 5 componentwise gives the claim. ■

Combining Theorem 1, Theorem 3 and Corollary 6 yields the complete statement: the L^2 state error at time t is bounded by $\sqrt{2t}$ times the sum of the empirical training residual, a quadrature term of order $(\log N_{col})^{d-1}/N_{col}$ with an explicitly monitored constant, and an approximation term of order $N^{-(k-1)}$ in the network width, plus a boundary contribution entering through the square root of its L^1 norm. Every term on the right is either measured after training or controlled by a design choice.

5.4. Where the Theory Stops

Corollary 6 requires $(h, q) \in H^k$ with $k \geq 3$, and Section 3.2 established that this fails. The closure law has a discontinuous derivative at $t = 0$ and at $t = T_c$; the resulting kinks propagate along the characteristics through those corners and reflect at both ends of the interval, so the solution is smooth only within the cells of the induced lattice and lies globally in $H^{3/2-\delta}$ at best. The approximation rate in Corollary 6 therefore degrades to $N^{-1/2+\delta}$ for a single global network, and no amount of training reaches the residual level that the smooth-solution theory would predict.

Two mitigations follow directly. Smoothing the closure law by a C^∞ mollification over a window short compared with $2L/a$ leaves the peak pressure essentially unchanged while removing the corner singularities at $t = 0$ and $t = T_c$; this is physically defensible, since a real guide-vane servo has finite jerk. Assigning independent subnetworks to the cells of the characteristic lattice restores smoothness within each cell and reduces the requirement to C^0 matching across interfaces. Section 4.6 implements the second and Section 7.7 measures its effect in the regime where it matters.

A second limitation is that Theorem 1 is an a posteriori estimate and says nothing about whether the optimiser will find a parameter vector with small residual. The nonconvexity of the training problem is not addressed here, and no statement in this section should be read as a convergence guarantee for Algorithm 1.

5.5. The Certificate

The practical content of the analysis is a post-training procedure. After Phase 2, evaluate $\|r_c\|_{L^2(\Omega_t)}$, $\|r_m\|_{L^2(\Omega_t)}$ and $\|r_v\|_{L^1(0,t)}$ on a fresh, denser Sobol point set that was not used in training, bound M by the range of the network output plus the range of the observed data, and evaluate the right-hand side of Theorem 1. The result is an upper bound on the state error that requires no knowledge of the exact solution and can therefore be computed online. Section 7.8 compares this bound with the error measured against the reference solution and quantifies its tightness, which is the only honest way to report a bound of this kind.

6. EXPERIMENTAL SETUP

6.1. Case Study

The test configuration represents a single unit of a high-head cascade in mountainous terrain, with a penstock geometry and closure schedule chosen to place the transient in the severe but physically admissible regime. All parameters are listed below.

Quantity	Symbol	Value
Penstock length	L	1200 m
Internal diameter	D	3.2 m
Cross-sectional area	A	8.0425 m ²
Wave speed	a	1000 m s ⁻¹
Darcy–Weisbach friction factor	f	0.012
Reservoir head	H_r	300 m
Initial velocity	V_0	2.5 m s ⁻¹
Initial discharge	Q_0	20.106 m ³ s ⁻¹
Steady friction loss over L	–	1.4335 m
Initial head at valve	H_0^v	298.5665 m
Closure time	T_c	3.0 s
Simulation horizon	T	10 s
Wave return period	$2L/a$	2.4 s
Full reflection period	$4L/a$	4.8 s
Joukowski head rise	ΔH_J	254.842 m
Nondimensional horizon	$\tilde{T}$	8.333
Nondimensional closure time	$\tilde{T}_c$	2.5
Nondimensional friction number	λ	5.625×10^{-3}
Allievi pipeline characteristic	$aV_0/(2gH_r)$	0.4247
Closure number	$aT_c/(2L)$	1.25

Two dimensionless groups characterise the regime. The Allievi characteristic of 0.42 places the installation in the range where the Joukowski rise is a substantial fraction of the static head, so the transient governs the design. The closure number of 1.25 exceeds unity, meaning the closure is slower than the wave return period and the pressure rise is attenuated below the Joukowski value; this is the intended operating case. Closure numbers below unity are examined in Section 7.7 as a stress test.

6.2. Reference Solver

Ground truth is generated by a method of characteristics solver with the Courant number fixed at unity, $\Delta t = \Delta x/a$, so that the characteristics pass exactly through grid nodes and no interpolation is required. Writing $B = a/(gA)$ and $R_f = f\Delta x/(2gDA^2)$, the compatibility relations reduce at each interior node to a pair of linear equations solved directly for the updated state, with the friction term treated by the standard first-order splitting. The upstream node imposes $H = H_r$ and the downstream node solves the quadratic obtained by substituting the orifice characteristic into the positive compatibility relation. Appendix B gives the complete discrete formulation.

This choice is deliberate. The most common criticism of the method of characteristics in the machine learning literature is that it is inaccurate, which is false at unit Courant number, and comparing a surrogate against a deliberately weakened reference is not informative. The solver used here is exact for the linear part of the system and its only discretisation error arises from the friction splitting, which for $\lambda = 5.6 \times 10^{-3}$ is negligible over the horizon considered. A related numerical-computing perspective on the implementation of such reference solvers, discussed in a general engineering-mathematics setting rather than for hydraulic transients specifically, is given by Kamalov and Leung [42].

6.3. Verification and Validation of the Reference

The reference is verified against the two analytical limits of the theory before any learning result is reported.

Test	Configuration	Analytical value	Reference solver	Relative deviation
Joukowski limit	instantaneous closure, frictionless	254.8420 m	254.8420 m	1.1×10^{-16}
Michaud–Allievi limit	linear discharge closure, $T_c = 3$ s, frictionless	203.8736 m	203.8736 m	1.4×10^{-16}

Both limits are recovered to machine precision, which is the expected behaviour of a unit-Courant-number scheme and confirms that the implementation is free of dispersion and of sign errors in the compatibility relations.

Grid independence is established on the primary case with the orifice closure law and friction active.

Spatial reaches N	Δx	Δt	Peak head at valve	Deviation from $N = 960$
120	10.000 m	10.0 ms	488.92544 m	4.4×10^{-8}
240	5.000 m	5.0 ms	488.92543 m	1.0×10^{-8}
480	2.500 m	2.5 ms	488.92542 m	2.1×10^{-9}
960	1.250 m	1.25 ms	488.92542 m	–

The peak is invariant to eight significant figures across a factor of eight in resolution, so the reference is grid converged and $N = 240$ is used throughout. The peak occurs at $t = 2.400$ s, exactly the wave return period, as the theory requires.

The primary case attains a maximum head of 488.925 m at the valve, a rise of 190.36 m above the initial value, against 502.44 m predicted by the Michaud–Allievi formula. The 2.7% discrepancy is not numerical error but a real effect of the orifice characteristic: the Michaud formula assumes discharge decreases linearly in time, whereas the square-root dependence on instantaneous head sustains the discharge as the pressure rises and redistributes the deceleration. Linearising the valve condition would therefore introduce an error comparable to the accuracy claimed for the surrogate, which is why the nonlinear condition is retained.

The minimum head at the valve is 122.040 m, well above the vapour pressure limit, so assumption A5 holds throughout and the single-phase model is valid. The head at the valve at successive return periods is 488.93 m at 2.4 s, 176.59 m at 4.8 s, 423.11 m at 7.2 s and 177.19 m at 9.6 s, exhibiting the weak attenuation expected at $\lambda = 5.6 \times 10^{-3}$.

6.4. Scenario Family and Data

The parametric study covers $T_c \in [1.5, 8.0]$ s, $Q_0 \in [12, 24]$ m³ s⁻¹, $a \in [900, 1100]$ m s⁻¹ and $f \in [0.008, 0.018]$. Reference trajectories are computed on a Sobol design of 512 parameter vectors for training and 128 held-out vectors for testing. The dependence of the peak head on closure time, computed on the reference solver, is as follows.

T_c [s]	Closure number	Peak head [m]	Minimum head [m]	Peak rise / ΔH_J
0.4	0.17	554.75	46.59	1.005
1.2	0.50	554.56	46.63	1.005
2.0	0.83	554.37	46.66	1.004
2.4	1.00	554.28	46.68	1.003
3.0	1.25	488.93	122.04	0.747
4.0	1.67	431.45	208.90	0.521
5.0	2.08	400.41	257.28	0.400
6.0	2.50	381.06	252.87	0.324
8.0	3.33	358.29	264.70	0.234

The saturation of the peak at the Joukowsky value for every closure number below unity is the classical result and provides a further validation of the reference. It also delimits the learning problem: the map from closure time to peak pressure is piecewise smooth with a kink at $T_c = 2L/a$, and a surrogate trained across this boundary must represent that kink.

Synthetic instrumentation consists of pressure transducers at $x/L \in \{0.25, 0.50, 0.75, 1.00\}$ and a discharge measurement at the valve, sampled at 100 Hz, giving 1000 samples per channel over the horizon. Measurements are corrupted by additive Gaussian noise with standard deviation expressed as a percentage of ΔH_J for head and of Q_0 for discharge, and quantised to 12 bits over a range of twice the signal span. The default noise level is 2%.

6.5. Baselines

Five baselines are evaluated on identical data.

Analytical. The linearised Allievi solution, obtained by dropping the friction term and linearising the valve condition. It is exact in the frictionless linear limit and serves to isolate the contribution of the nonlinearities.

Data-driven multilayer perceptron. The architecture of Section 4 trained on $\mathcal{L}_d$ alone, with no residual, valve or causal terms and without the constrained ansatz. This isolates the contribution of the physics.

Gated recurrent network. Two layers of 256 units mapping the boundary schedule and past measurements to the state at the sensor locations, with linear interpolation in space. This represents the class of sequence surrogates in operational use.

DeepONet. Branch and trunk networks of five layers and 128 units, trained on the parametric family with the closure schedule as the branch input.

Fourier neural operator. A one-dimensional operator in space with 16 retained modes and four layers, advanced in time by a fixed stride.

6.6. Metrics

Accuracy is reported as the relative L_2 error over the full space-time domain,

$$\mathcal{R}_2(y) = \frac{(\sum_i (y_{pred,i} - y_{true,i})^2)^{1/2}}{(\sum_i y_{true,i}^2)^{1/2}},$$

evaluated on the reference grid with $N = 240$ and 2000 time levels, separately for head and discharge. Because operational risk is dominated by extrema rather than by mean behaviour we additionally report the signed relative error in the peak head, the absolute error in the time of the peak, the signed relative error in the minimum head, the phase lag estimated by cross-correlation of the valve head trace against the reference over the full horizon, and the residual norms entering the certificate of Section 5.5.

6.7. Protocol

Every learning result is reported as the mean and standard deviation over 10 independent runs differing in network initialisation, Fourier matrix draw and Sobol scramble. Comparisons between methods use the Wilcoxon signed-rank test over paired runs, and differences are described as significant only at $p < 0.01$. Hyperparameters are fixed across all experiments at the values in Appendix D and were selected on a validation scenario with $T_c = 3.5$ s that appears in no reported result. Training used a single NVIDIA A100 with 40 GB of memory; Phase 1 runs in single precision and Phase 2 in double precision, which is necessary for L-BFGS to make progress once the gradient norm falls below approximately 10^{-6} .

6.8. Reproducibility

The reference solver, the verification suite of Section 6.3, the grid independence study and the closure-time sweep are fully specified by Appendix B and the parameter table of Section 6.1, and every number reported in those sections was regenerated for this manuscript. Learning results depend on the training procedure of Algorithm 1 and on the

hardware described above; the certificate of Section 5.5 is reported alongside the measured error precisely so that a reader can check the internal consistency of the learning results without rerunning them.

7. RESULTS

7.1. Accuracy on the Primary Case

The table reports relative L_2 errors over the full space-time domain for the primary case, as mean and standard deviation over 10 runs.

Method	$\mathcal{R}_2(H)$ [%]	$\mathcal{R}_2(Q)$ [%]
Linearised Allievi solution	7.83	9.12
Data-driven multilayer perceptron	18.74 ± 2.41	21.35 ± 3.06
Gated recurrent network	12.38 ± 1.77	14.92 ± 2.24
DeepONet	4.62 ± 0.51	5.38 ± 0.63
Fourier neural operator	2.91 ± 0.34	3.44 ± 0.41
Proposed model	1.14 ± 0.12	1.32 ± 0.15

The ordering is stable across all 10 runs and every pairwise comparison against the proposed model is significant at $p < 0.01$. Three observations are worth separating from the ranking itself.

The linearised analytical solution, at 7.83%, already outperforms both purely data-driven surrogates. A model that cannot beat the closed-form solution of the linearised problem has learned nothing about the system beyond its gross shape, and the fact that this comparison is rarely reported in the applied literature is a weakness of that literature rather than a favourable result for the models in question.

The gap between the operator-learning baselines and the proposed model is a factor of two to four, considerably smaller than the gap to the unstructured baselines. Operator methods impose no physics but do impose an inductive bias suited to the smooth, translation-covariant structure of wave propagation, and this recovers most of the benefit on the forward problem. Their disadvantage appears in the extremal metrics of Section 7.2 and in the noise and sparsity studies, where the absence of a residual leaves nothing to regularise the fit.

The residual of the trained model is not zero, and the certificate of Section 5.5 correctly reports this. The measured error of 1.14% should be read together with the certificate value of 13.4% in Section 7.8: the model is accurate, and the guarantee is much weaker than the accuracy.

7.2. Extremal and Phase Accuracy

Mean absolute errors over 10 runs are reported. Peak and minimum errors are signed relative errors in the head at the valve; phase lag is estimated by cross-correlation of the valve head trace over the full horizon.

Method	Peak head error [%]	Peak timing error [ms]	Minimum head error [%]	Phase lag [ms]
Data-driven multilayer perceptron	−16.8	148	+22.4	380
Gated recurrent network	−9.4	87	+13.1	214
DeepONet	−4.1	41	+5.9	96
Fourier neural operator	−2.71	34	+3.8	61
Proposed model	−0.63	11	+1.42	8

Every baseline underestimates the peak and overestimates the minimum, that is, all of them damp the transient. This is the characteristic failure of a regression fit to an oscillatory signal, and it is the failure mode with the worst operational consequences, since a state estimator that systematically underestimates pressure extrema is unsafe in exactly the situation for which it was installed. The proposed model exhibits the same bias but reduced by more than an order of magnitude relative to the recurrent baseline, and the residual bias is traceable to the spectral content of the front rather than to the training objective. In absolute terms the peak head error of 0.63% corresponds to 3.08 m on a rise of 190.36 m, and the timing error of 11 ms is small against the 2.4 s return period but is not negligible against a 10 ms control cycle.

7.3. Robustness to Measurement Noise

Noise is applied to all measurement channels at the stated level, with the residual and valve terms unchanged.

Noise level [% of range]	Proposed model $\mathcal{R}_2(H)$ [%]	Data-driven baseline $\mathcal{R}_2(H)$ [%]
0	1.09 ± 0.11	18.74 ± 2.41
1	1.11 ± 0.11	19.62 ± 2.55
2	1.14 ± 0.12	20.88 ± 2.74
5	1.42 ± 0.18	24.61 ± 3.42
10	2.37 ± 0.41	33.07 ± 4.88

Degradation is sublinear in the noise level up to 5% and becomes appreciable only at 10%. The mechanism is visible in the residual norms, which vary by less than 8% across the entire sweep: the residual term does not fit the noise because noise is not a solution of the governing system, so the effective capacity available for overfitting is small. This is the clearest empirical benefit of the formulation and it is a direct consequence of the constraint, not of any regularisation hyperparameter.

7.4. Sensor Sparsity

Channels are removed in order of increasing distance from the valve, retaining the valve discharge measurement last.

Measurement channels	$\mathcal{R}_2(H)$ [%]	$\mathcal{R}_2(Q)$ [%]
5, all	1.14 ± 0.12	1.32 ± 0.15
4	1.19 ± 0.13	1.38 ± 0.16
3	1.29 ± 0.15	1.51 ± 0.19
2	1.41 ± 0.18	1.66 ± 0.22
1, valve discharge only	1.58 ± 0.24	1.84 ± 0.29
0, pure forward solve	1.63 ± 0.26	1.91 ± 0.31

The error rises by only 50% as the instrumentation is reduced from five channels to none, because the boundary conditions and the residual already determine the solution; measurements refine an estimate that is well posed without them. The corresponding sweep for the data-driven baseline is not meaningful, since with no measurements it has no training signal at all. The practically relevant reading is that a single existing transducer is sufficient for the estimator to track a real installation, which matters because retrofitting instrumentation on a pressurised penstock is expensive and often deferred indefinitely.

7.5. Parametric Generalisation

The parametric model of Section 4.7 is trained on 512 Sobol parameter vectors and evaluated on 128 held-out vectors.

Evaluation set	$\mathcal{R}_2(H)$ [%] mean	95th percentile	Worst case
Held-out, full range	1.87 ± 0.44	3.16	4.12
Held-out, closure number > 1.2	1.54 ± 0.31	2.44	2.98
Held-out, closure number $\in [0.9, 1.2]$	3.28 ± 0.71	4.02	4.12
Extrapolation, $T_c = 1.0$ s	6.91 ± 1.24	–	–

Accuracy is uniform over the family except in a neighbourhood of the closure number equal to unity, where the error roughly doubles. This is the kink identified in Section 6.4: the map from closure time to peak pressure is not differentiable at $T_c = 2L/a$, and a smooth network must smear it. The failure is therefore predicted by the structure of the problem rather than discovered empirically, and it can be removed by partitioning the parameter domain at the kink, at the cost of maintaining two models. Extrapolation beyond the training range degrades sharply and the model should not be used there; this limitation is shared by all the operator baselines and is not specific to the physics-informed formulation, since the residual constrains the solution given the parameters but does not extend the parameter range over which the network has been fitted.

The peak head across the held-out family is predicted with a mean absolute relative error of 1.03% and a worst case of 3.4%, the latter again at the closure number kink.

7.6. Ablations

All entries are for the primary case with 2% noise.

Configuration	$\mathcal{R}_2(H)$ [%]	Change
Full model	1.14 ± 0.12	–
Without causal weighting	3.87 ± 0.92	$\times 3.4$
Without neural tangent kernel weighting	2.65 ± 0.58	$\times 2.3$
Without Fourier embedding	5.12 ± 1.13	$\times 4.5$
Without exact initial and boundary imposition	2.08 ± 0.31	$\times 1.8$
Without L-BFGS phase	4.41 ± 0.47	$\times 3.9$
Latin hypercube instead of Sobol collocation	1.31 ± 0.15	$\times 1.15$
Sine activation instead of tanh	1.21 ± 0.19	not significant
Rectified linear activation	12.83 ± 3.14	$\times 11.3$
Without residual term, data only	18.74 ± 2.41	$\times 16.4$

The Fourier embedding and the causal weighting are the two components whose removal is most damaging, and both address the same underlying difficulty: the solution is oscillatory in time, and an unmodified network fits the low-frequency envelope first and the reflection structure late or never. The exact imposition of initial and boundary data contributes a factor of 1.8, consistent in direction with Proposition 2 though the theory predicts only that the constant deteriorates, not by how much. The rectified linear activation fails as expected, since its derivative is piecewise constant and the residual cannot be driven to zero except in a measure-theoretic sense that the collocation estimator does not see.

The scale of the Fourier embedding behaves as the spectral argument of Section 4.2 predicts.

σ_F	1	2	3	5	10
$\mathcal{R}_2(H)$ [%]	3.24	1.22	1.14	1.19	2.87

Performance is flat between 2 and 5 and degrades on both sides, at low σ_F through the reappearance of spectral bias and at high σ_F through a high-frequency ripple superposed on a correct envelope.

7.7. The Fast-Closure Regime

The primary case has a closure number of 1.25 and therefore a comparatively mild front. Reducing the closure time to $T_c = 0.4$ s, a closure number of 0.17, produces the full Joukowski rise and a much steeper wavefront: measured on the reference solution, the maximum rate of head change at the valve rises from 203.5 m s^{-1} to 1719.2 m s^{-1} , a factor of 8.4, and the total variation of the valve head trace over the horizon rises from 1250.5 m to 2271.7 m.

Configuration	$\mathcal{R}_2(H)$ [%]	Peak head error [%]
Single global network	8.41 ± 1.36	–4.83
Mollified closure law	6.12 ± 0.94	–3.41
Characteristic-aligned decomposition, 7 cells	2.63 ± 0.35	–1.12
Both	2.41 ± 0.31	–0.97
Reference solver	–	–

A single network loses roughly a factor of seven in accuracy relative to the primary case, and the peak error of nearly 5% would be unacceptable in a protection-adjacent application. The characteristic decomposition recovers most of the loss, which supports the diagnosis of Section 5.4 that the limitation is the global smoothness assumption rather than capacity or optimisation. It does not recover all of it, and at closure numbers well below unity the method of characteristics remains the better tool. We regard this as the principal accuracy limitation of the approach and state it plainly: the surrogate is appropriate for the controlled closure schedules of normal and emergency operation and is not appropriate for the near-instantaneous events associated with valve slam or protection trips.

7.8. Tightness of the Certificate

The certificate of Section 5.5 was evaluated on a fresh Sobol set of 4×10^5 points for the primary case.

Quantity	Value
$\ r_c\ _{L^2(\Omega)}$	4.7×10^{-3}
$\ r_m\ _{L^2(\Omega)}$	6.1×10^{-3}
$\ r_v\ _{L^1(0,\tilde{T})}$	2.3×10^{-3}
M	6.0×10^{-2}
Bound of Theorem 1 at $t = \tilde{T}$	2.09×10^{-3}
Implied bound on $(\ e\ ^2 + \ \varepsilon\ ^2)^{1/2}$	4.57×10^{-2}
Implied relative error bound	13.4%
Measured relative error at $t = \tilde{T}$	1.21%
Overestimation factor	11.1

The certificate is valid and loose by an order of magnitude, which is the expected behaviour of an energy estimate that discards the friction dissipation entirely and bounds the boundary contribution by a supremum. It is nevertheless useful: it is computable without the exact solution, it is monotone in the residuals, and an order of magnitude of headroom is a workable safety factor for an operational alarm threshold. We note that the two terms contribute comparably, 9.9×10^{-4} from the interior residuals and 1.1×10^{-3} from the valve term, which again indicates that a hard imposition of the valve condition would be the most valuable remaining improvement.

The quadrature term of Theorem 3 was also evaluated. With $N_{col} = 4 \times 10^4$, $d = 2$ and a monitored Hardy–Krause variation of approximately 35, the term contributes 9.3×10^{-3} , which exceeds the empirical residual itself. Increasing N_{col} tenfold reduces the quadrature term by a factor of nine and increases training time by a factor of 3.1 while leaving the measured error unchanged at $1.13 \pm 0.12\%$. The quadrature bound is therefore not binding in practice at these sample counts and its apparent dominance reflects the looseness of the variation estimate rather than a genuine sampling deficiency.

7.9. Computational Cost

This section corrects a claim that recurs in the applied literature and appeared in the first draft of this work. All timings are wall-clock, for the primary case with $N = 240$ and a 10 s horizon, and the reference solver timings were measured rather than cited.

Task	Method	Time
Full-field reconstruction, 241×2001 states	Reference solver, vectorised array implementation	33 ms
Full-field reconstruction	Reference solver, scalar loop implementation	118 ms
Full-field reconstruction	Proposed model, single forward pass	130 ms
Pointwise query, batch of 32 states	Reference solver	33 ms
Pointwise query, batch of 32 states	Proposed model	0.9 ms
50 candidate closure schedules, 8 sensors, 20 horizon steps	Reference solver	1650 ms
50 candidate closure schedules, 8 sensors, 20 horizon steps	Proposed model, batched	1.4 ms
Online parameter assimilation, 12 Gauss–Newton iterations	Proposed model	6.4 ms
Offline training, single scenario	Proposed model	38 min
Offline training, parametric family	Proposed model	4.2 h

For full-field reconstruction the proposed model is approximately four times slower than a competently vectorised reference solver, not two orders of magnitude faster. Reports of speedups in the range of 98% are obtained by comparing a graphics-processor surrogate against a scalar-loop reference implementation, and they measure the implementation rather than the method. We state the corrected comparison because the opposite claim, once propagated, leads to design decisions that the evidence does not support.

The genuine advantages are structural and are visible in the remaining rows. The reference solver must integrate from the initial condition to reach any state, so the cost of a single pointwise query equals the cost of the entire trajectory, whereas the surrogate evaluates in constant time at any point. The gap widens by three orders of magnitude in the predictive control setting, where many candidate schedules must be scored over a short horizon and the surrogate evaluates all of them in one batched pass. Assimilation of new telemetry costs 6.4 ms and requires no retraining. The surrogate is also differentiable with respect to time and to the operating parameters at no additional cost, whereas obtaining the same sensitivities from the reference solver requires an adjoint implementation. Finally, the deployed

model occupies 0.40 MB of parameters against 3.8 MB for a single stored reference field, which matters on embedded governor hardware.

The correct summary is therefore that the surrogate does not accelerate simulation but changes the cost model of state estimation from per-trajectory to per-query, and that this is what makes it useful in a control loop.

7.10. Inverse Identification

Treating a and f as trainable parameters on the primary case with 2% noise and five measurement channels gives the following.

Parameter	True value	Estimate	Relative error
Wave speed a	1000 m s ⁻¹	996.4 ± 4.1 m s ⁻¹	0.36%
Friction factor f	0.012	0.01243 ± 0.00068	3.58%

The asymmetry is informative and physically expected. The wave speed sets the return period $2L/a$, which is directly observable in the timing of the pressure peaks and is therefore strongly identified from a short record. The friction factor acts only through the attenuation between successive peaks, which at $\lambda = 5.6 \times 10^{-3}$ amounts to a few per cent per period, so a 10 s record contains little information about it. Extending the record to 20 s reduces the friction error to 1.41% while leaving the wave speed error unchanged at 0.34%, confirming the diagnosis. Practitioners should treat identified friction factors from short transient records with corresponding caution, whatever method produced them.

8. DISCUSSION

8.1. What the Evidence Supports

Three claims are supported by the results. First, embedding the residual of the water hammer system in the training objective reduces the error on the primary case by a factor of sixteen relative to the same architecture trained on data alone, and by a factor of two and a half relative to the strongest operator-learning baseline, with the largest gains on the extremal and phase metrics that determine operational value. Second, the resulting estimator is robust to the noise and sparsity characteristic of penstock instrumentation, degrading by about 30% under 5% measurement noise and by 50% when all instrumentation is removed. Third, the error is controlled by a computable certificate whose constant grows as the square root of the horizon, which makes residual-based monitoring meaningful rather than vacuous over the multi-period horizons of interest.

Two claims commonly made for this class of method are not supported. The surrogate is not faster than a competent classical solver at reconstructing a full field; it is approximately four times slower. And the physics term does not confer extrapolation in the parameter domain: accuracy outside the trained range of closure times degrades to 6.9%, comparable to the operator baselines. The residual constrains which functions of (x, t) are admissible for a given parameter vector, and it says nothing about parameter vectors the network has never been fitted at.

8.2. Modelling Limitations

The quasi-steady friction assumption A4 is the most consequential simplification. Unsteady friction models attribute additional attenuation to the wall shear associated with the transient velocity profile, and the discrepancy is known to be significant after the first few periods. The convolution-type models of Zielke and of Vardy and Brown are non-local in time and therefore not directly expressible as a pointwise residual, which is a genuine structural obstacle for the present formulation rather than an oversight; the instantaneous-acceleration models of the Brunone family are local and could be incorporated without modification, at the cost of an additional identified parameter that Section 7.10 suggests would be poorly determined from short records. Because the reference solver makes the same assumption, the reported errors measure agreement with the model, not with a physical penstock, and validation against instrumented field transients remains necessary before any operational claim.

Assumption A5 excludes column separation. The primary case has a minimum head of 122 m and is comfortably within the single-phase regime, but the fast-closure cases of Section 7.7 reach 46.6 m, and a configuration with a higher initial velocity or a longer conduit would cross into cavitation. Discrete vapour cavity models introduce genuine discontinuities in the state, and a smooth network is the wrong representation for them.

The study treats a single conduit with a reservoir and a valve. Real cascades contain surge tanks, manifolds, multiple units and tailrace tunnels, and the junction conditions coupling them are algebraic constraints linking the states of several conduits. Nothing in the formulation forbids this, and Section 8.5 sketches the extension, but it has not been demonstrated here.

8.3. Method Limitations

The training problem is nonconvex and no result in Section 5 asserts that Algorithm 1 will find a parameter vector with small residual. The 10-run standard deviations of 10 to 30% of the mean error indicate that the optimiser does

not converge to equivalent solutions, and the ablation without the L-BFGS phase shows that the outcome depends materially on the optimiser schedule. This is the least satisfactory aspect of the method and it is not specific to this application.

The offline training cost of 4.2 h for the parametric family is charged once but must be recharged whenever the plant configuration changes materially. Section 4.8 mitigates this by exposing the slowly drifting parameters at inference time, but a change in conduit geometry or the installation of a surge tank requires retraining.

The certificate of Section 5.5 is loose by a factor of 11. It is adequate for setting a conservative alarm threshold and inadequate as a tight accuracy guarantee. The dominant looseness comes from discarding the friction dissipation and from bounding the boundary contribution by a supremum; a sharper estimate that retains the dissipation would require a lower bound on $|\varepsilon|$ over the boundary, which is not available a posteriori.

8.4. Deployment Considerations

An operational deployment should treat the surrogate as an estimator running alongside, not in place of, a classical solver. The residual evaluated on live telemetry is the natural monitor: it is computable at inference cost, it is zero for a correct model applied to a correct plant, and it rises under either model error or plant change. A practical architecture therefore evaluates the certificate continuously, uses the surrogate inside the control loop while the certificate remains below a threshold set from Section 7.8, and falls back to a conservative envelope or to an offline solver run when it does not. The residual monitor additionally functions as an anomaly detector, since air entrainment, partial blockage or valve malfunction all manifest as a persistent rise in the residual that no amount of data fitting can absorb.

The regime boundary established in Section 7.7 should be enforced explicitly. Closure numbers below unity produce fronts the surrogate resolves poorly, and since the closure schedule is commanded rather than observed, this condition is known in advance and the fallback can be triggered before the event rather than after.

8.5. Extension to Cascades

Scaling to a multi-unit cascade requires composing local estimators across a network of conduits joined at manifolds, surge tanks and free surfaces. The structure of the composition is dictated by the energy estimate of Section 5.1 rather than by the learning architecture. The proof integrates a local energy identity over the spatial interval and the coupling between subdomains enters only through the boundary flux term $[e\varepsilon]$. If the junction conditions are such that the sum of the outgoing fluxes over all conduits meeting at a junction is non-positive, which is the discrete statement that a junction does not create energy, then the local estimates add and a certificate for the cascade follows from the certificates of its parts with no additional constant. Establishing this for the standard junction models is a concrete and tractable objective.

This is also where the categorical framing suggested in the original draft acquires content. An indexed family of local state spaces over a diagram of conduits and junctions, assembled by the Grothendieck construction into a total space, is a precise way of saying that global states are compatible families of local states; the compositionality property described above is then the statement that the certificate is a lax monoidal functor on that structure. The value of the formalism, if any, lies in making the compositional guarantee provable once for a class of junction models rather than case by case, and we regard it as a hypothesis to be tested rather than a result.

Explainability is a second direction with a concrete target. The trained network is differentiable in x , t and μ , so sensitivities are available directly, and the useful object for a grid operator is not a saliency map over inputs but the derivative of the predicted peak pressure with respect to the commanded closure schedule. That quantity is a single backward pass and is immediately actionable, which is a stronger argument than the generic case for post hoc attribution.

9. CONCLUSION

This paper has developed and analysed a physics-informed neural network for state estimation of hydraulic transients in high-head penstocks. The analysis exploits two structural properties of the water hammer system, the symmetry of its flux matrix and the monotonicity of its quadratic friction term, to obtain an a posteriori error estimate whose constant grows as the square root of the horizon rather than exponentially, and combines it with quadrature and approximation bounds into a certificate expressed entirely in computable quantities. The analysis also determines two design choices: boundary and initial conditions are imposed exactly because boundary residuals enter the estimate through a weaker norm than interior residuals, and the space-time domain is decomposed along characteristics because the solution is piecewise rather than globally smooth. Both choices were validated empirically, contributing factors of 1.8 and 3.2 respectively in the regimes where they apply.

Empirically, the estimator attains 1.14% relative error in head and 1.32% in discharge against a reference solver verified to machine precision against the Joukowski and Michaud–Allievi limits, predicts the peak pressure to within 0.63% and its timing to within 11 ms, degrades by less than a third under 5% measurement noise, retains 1.63% accuracy with no instrumentation at all, and recovers the acoustic wave speed to 0.36% in the inverse setting. A parametric formulation amortises training across the operating envelope and supports assimilation of new telemetry in 6.4 ms.

The cost accounting has been corrected. Against a competently implemented reference solver the surrogate is slower at full-field reconstruction, and the frequently quoted order-of-magnitude speedups measure implementation quality rather than method. What the surrogate changes is the cost model: from per-trajectory to per-query, from re-simulation to re-conditioning, and from a black-box solver to a differentiable one. In a predictive control loop scoring fifty candidate closure schedules this is a three-order-of-magnitude difference, and it is the reason the approach is worth deploying.

The limitations are equally clear. Fast closures with closure numbers below unity produce fronts that a smooth network resolves poorly even with characteristic decomposition, extrapolation outside the trained parameter range is unreliable, the training problem is nonconvex with run-to-run variation of 10 to 30%, and the certificate is loose by an order of magnitude. Each of these is a specific and testable deficiency rather than a general caveat, and the immediate work is to address the first by hard imposition of the valve characteristic and hybrid treatment of the front, and the second by coupling the estimator to a classical solver that supplies training data on demand near the boundary of the trained envelope.

REFERENCES

- [1] E. B. Wylie and V. L. Streeter. *Fluid Transients in Systems*. Prentice Hall, 1993.
- [2] M. H. Chaudhry. *Applied Hydraulic Transients*. Springer, 3rd edition, 2014.
- [3] M. S. Ghidaoui, M. Zhao, D. A. McInnis, and D. H. Axworthy. A review of water hammer theory and practice. *Applied Mechanics Reviews*, 58(1):49–76, 2005.
- [4] W. Zielke. Frequency-dependent friction in transient pipe flow. *Journal of Basic Engineering*, 90(1):109–115, 1968.
- [5] A. E. Vardy and J. M. B. Brown. Transient turbulent friction in smooth pipe flows. *Journal of Sound and Vibration*, 259(5):1011–1036, 2003.
- [6] A. Bergant, A. R. Simpson, and J. Vitkovsky. Developments in unsteady pipe flow friction modelling. *Journal of Hydraulic Research*, 39(3):249–257, 2001.
- [7] M. Raissi, P. Perdikaris, and G. E. Karniadakis. Physics-informed neural networks: a deep learning framework for solving forward and inverse problems involving nonlinear partial differential equations. *Journal of Computational Physics*, 378:686–707, 2019.
- [8] J. Sirignano and K. Spiliopoulos. DGM: a deep learning algorithm for solving partial differential equations. *Journal of Computational Physics*, 375:1339–1364, 2018.
- [9] G. E. Karniadakis, I. G. Kevrekidis, L. Lu, P. Perdikaris, S. Wang, and L. Yang. Physics-informed machine learning. *Nature Reviews Physics*, 3:422–440, 2021.
- [10] S. Cuomo, V. S. di Cola, F. Giampaolo, G. Rozza, M. Raissi, and F. Piccialli. Scientific machine learning through physics-informed neural networks: where we are and what’s next. *Journal of Scientific Computing*, 92:88, 2022.
- [11] S. Wang, Y. Teng, and P. Perdikaris. Understanding and mitigating gradient flow pathologies in physics-informed neural networks. *SIAM Journal on Scientific Computing*, 43(5):A3055–A3081, 2021.
- [12] S. Wang, X. Yu, and P. Perdikaris. When and why PINNs fail to train: a neural tangent kernel perspective. *Journal of Computational Physics*, 449:110768, 2022.
- [13] A. S. Krishnapriyan, A. Gholami, S. Zhe, R. M. Kirby, and M. W. Mahoney. Characterizing possible failure modes in physics-informed neural networks. In *Advances in Neural Information Processing Systems*, 2021.
- [14] S. Wang, S. Sankaran, and P. Perdikaris. Respecting causality for training physics-informed neural networks. *Computer Methods in Applied Mechanics and Engineering*, 421:116813, 2024.
- [15] N. Rahaman, A. Baratin, D. Arpit, F. Draxler, M. Lin, F. Hamprecht, Y. Bengio, and A. Courville. On the spectral bias of neural networks. In *International Conference on Machine Learning*, 2019.
- [16] M. Tancik, P. P. Srinivasan, B. Mildenhall, S. Fridovich-Keil, N. Raghavan, U. Singhal, R. Ramamoorthi, J. T. Barron, and R. Ng. Fourier features let networks learn high frequency functions in low dimensional domains. In *Advances in Neural Information Processing Systems*, 2020.
- [17] A. D. Jagtap, E. Kharazmi, and G. E. Karniadakis. Conservative physics-informed neural networks on discrete domains for conservation laws. *Computer Methods in Applied Mechanics and Engineering*, 365:113028, 2020.

- [18] Y. Shin, J. Darbon, and G. E. Karniadakis. On the convergence of physics informed neural networks for linear second-order elliptic and parabolic type PDEs. *Communications in Computational Physics*, 28(5):2042–2074, 2020.
- [19] S. Mishra and R. Molinaro. Estimates on the generalization error of physics-informed neural networks for approximating PDEs. *IMA Journal of Numerical Analysis*, 42(2):981–1022, 2022.
- [20] T. De Ryck, S. Lanthaler, and S. Mishra. On the approximation of functions by tanh neural networks. *Neural Networks*, 143:732–750, 2021.
- [21] T. De Ryck and S. Mishra. Error analysis for physics-informed neural networks approximating a class of inverse problems for PDEs. *Advances in Computational Mathematics*, 48:79, 2022.
- [22] P. D. Lax. *Hyperbolic Systems of Conservation Laws and the Mathematical Theory of Shock Waves*. SIAM, 1973.
- [23] C. M. Dafermos. *Hyperbolic Conservation Laws in Continuum Physics*. Springer, 4th edition, 2016.
- [24] L. Lu, P. Jin, G. Pang, Z. Zhang, and G. E. Karniadakis. Learning nonlinear operators via DeepONet based on the universal approximation theorem of operators. *Nature Machine Intelligence*, 3:218–229, 2021.
- [25] Z. Li, N. Kovachki, K. Azizzadenesheli, B. Liu, K. Bhattacharya, A. Stuart, and A. Anandkumar. Fourier neural operator for parametric partial differential equations. In *International Conference on Learning Representations*, 2021.
- [26] G. S. Misyris, A. Venzke, and S. Chatzivasileiadis. Physics-informed neural networks for power systems. In *IEEE Power and Energy Society General Meeting*, 2020.
- [27] J. Stiasny, S. Chatzivasileiadis, and G. S. Misyris. Physics-informed neural networks for non-linear system identification for power system dynamics. In *IEEE PowerTech*, 2021.
- [28] G. Kissas, Y. Yang, E. Hwuang, W. R. Witschey, J. A. Detre, and P. Perdikaris. Machine learning in cardiovascular flows modeling: predicting arterial blood pressure from non-invasive 4D flow MRI data using physics-informed neural networks. *Computer Methods in Applied Mechanics and Engineering*, 358:112623, 2020.
- [29] H. Niederreiter. *Random Number Generation and Quasi-Monte Carlo Methods*. SIAM, 1992.
- [30] A. B. Owen. Scrambled net variance for integrals of smooth functions. *Annals of Statistics*, 25(4):1541–1562, 1997.
- [31] R. H. Byrd, P. Lu, J. Nocedal, and C. Zhu. A limited memory algorithm for bound constrained optimization. *SIAM Journal on Scientific Computing*, 16(5):1190–1208, 1995.
- [32] D. P. Kingma and J. Ba. Adam: a method for stochastic optimization. In *International Conference on Learning Representations*, 2015.
- [33] H.-O. Kreiss. Initial boundary value problems for hyperbolic systems. *Communications on Pure and Applied Mathematics*, 23(3):277–298, 1970.
- [34] D. I. Spivak. *Category Theory for the Sciences*. MIT Press, 2014.
- [35] B. Padhy, P. S. Mahapatra, and T. Samantara. Adaptive hybrid B-spline PINN for fractional equations. *Annals of Mathematics and Computer Science*, 35, 2026. DOI: 10.56947/amcs.v35.864.
- [36] S. Zongo and B. A. Kyelem. 3D Navier–Stokes equation with multivalued friction. *Annals of Mathematics and Computer Science*, 35, 2026. DOI: 10.56947/amcs.v35.873.
- [37] R. Obogi and M. N. Evans. Spectral gaps and approximate invariance in Hilbert spaces. *Annals of Mathematics and Computer Science*, 30:85–92, 2025. DOI: 10.56947/amcs.v30.650.
- [38] Z. Souleymane, B. Lamien, M. Beidari, and T. Inoussa. Numerical simulation of pollutants transport in ground-water using deep neural networks informed by physics. *Gulf Journal of Mathematics*, 16(2):337–352, 2024. DOI: 10.56947/gjom.v16i2.1863.
- [39] S. Sore, B. Leye, and Y. Simporé. A Godunov-type scheme for the Saint-Venant–Exner equations with a moving steady state. *Gulf Journal of Mathematics*, 17(2):166–189, 2024. DOI: 10.56947/gjom.v17i2.1937.

- [40] A. Fadhil and A. Hassan. Approximation by new Baskakov operators. *Gulf Journal of Mathematics*, 19(2):442–450, 2025. DOI: 10.56947/gjom.v19i2.2687.
- [41] F. Kamalov and H. H. Leung. Fixed point theory: a review. *arXiv preprint arXiv:2309.03226*, 2023. DOI: 10.48550/arXiv.2309.03226.
- [42] F. Kamalov and H. Leung. Numerical computing in engineering mathematics. In *2022 Advances in Science and Engineering Technology International Conferences (ASET)*, pages 1–5, 2022. DOI: 10.1109/aset53988.2022.9734960.
- [43] F. Kamalov, S. Moussa, and J. A. Reyes. Regularized information loss for improved model selection. In *Intelligent Communication Technologies and Virtual Mobile Networks*, pages 801–811, 2023. DOI: 10.1007/978-981-99-1767-9_58.

APPENDIX A. DERIVATION OF THE GOVERNING SYSTEM

Consider a control volume of length δx in a prismatic conduit of area A inclined at angle α , containing a fluid of density ρ and bulk modulus K in a wall of thickness e and Young's modulus E .

Continuity. Conservation of mass gives $\partial_t(\rho A) + \partial_x(\rho AV) = 0$. Expanding and dividing by ρA ,

$$\frac{1}{\rho} \frac{d\rho}{dt} + \frac{1}{A} \frac{dA}{dt} + \frac{\partial V}{\partial x} = 0.$$

The first term is $\frac{1}{K} \frac{dp}{dt}$ by definition of the bulk modulus. The second follows from the elastic response of the wall: for a thin-walled conduit the hoop strain gives $\frac{1}{A} \frac{dA}{dt} = \frac{Dc_1}{Ee} \frac{dp}{dt}$, where c_1 depends on the axial restraint. Combining,

$$\frac{1}{\rho a^2} \frac{dp}{dt} + \frac{\partial V}{\partial x} = 0, \quad a^2 = \frac{K/\rho}{1 + \frac{KDc_1}{Ee}}.$$

Writing $p = \rho g(H - z)$ and neglecting the convective contribution to the material derivative under A3 gives the continuity equation of Section 3.1 after multiplication by a^2/g and substitution of $Q = AV$.

Momentum. Newton's second law applied to the same control volume gives

$$\rho \frac{dV}{dt} = -\frac{\partial p}{\partial x} - \rho g \sin \alpha - \frac{4\tau_w}{D},$$

with wall shear $\tau_w = \rho f V |V|/8$ under A4. Substituting $p = \rho g(H - z)$, noting that $\partial_x z = \sin \alpha$ cancels the gravitational term, neglecting convective acceleration and multiplying by A/ρ gives the momentum equation of Section 3.1.

Consistency of the case study. The configuration of Section 6.1 is internally consistent as a physical design. With $K = 2.19$ GPa, $\rho = 999$ kg m⁻³, $E = 210$ GPa and $c_1 = 1$, the specified wave speed of 1000 m s⁻¹ implies a wall thickness of 28.0 mm for $D = 3.2$ m. The corresponding hoop stress is 168 MPa under the static head of 300 m and 274 MPa at the transient peak of 488.93 m, which is within the yield strength of penstock-grade steel but leaves a modest margin. The case study is therefore a severe but admissible design point, which is the intended regime.

APPENDIX B. REFERENCE SOLVER

The conduit is divided into N reaches of length $\Delta x = L/N$ with $\Delta t = \Delta x/a$, so the Courant number is unity and the characteristics pass through grid nodes. Define

$$B = \frac{a}{gA}, \quad R_f = \frac{f\Delta x}{2gDA^2}.$$

At an interior node i at time level $n+1$ the compatibility relations along the two characteristics reaching it are

$$\begin{aligned} C^+ : \quad H_i^{n+1} &= C_P - B_P Q_i^{n+1}, & C_P &= H_{i-1}^n + B Q_{i-1}^n, & B_P &= B + R_f |Q_{i-1}^n|, \\ C^- : \quad H_i^{n+1} &= C_M + B_M Q_i^{n+1}, & C_M &= H_{i+1}^n - B Q_{i+1}^n, & B_M &= B + R_f |Q_{i+1}^n|, \end{aligned}$$

whence $Q_i^{n+1} = (C_P - C_M)/(B_P + B_M)$ and $H_i^{n+1} = C_P - B_P Q_i^{n+1}$.

At the upstream node only C^- is available and is closed by $H_0^{n+1} = H_r$, giving $Q_0^{n+1} = (H_r - C_M)/B_M$ with C_M, B_M formed from node 1.

At the downstream node only C^+ is available and is closed by the orifice characteristic. Writing $C_v = (\tau^{n+1} Q_0)^2/H_0^v$ and eliminating H_N^{n+1} gives the quadratic $(Q_N^{n+1})^2 + C_v B_P Q_N^{n+1} - C_v C_P = 0$ with the physically admissible root

$$Q_N^{n+1} = \frac{1}{2} \left(-C_v B_P + \sqrt{(C_v B_P)^2 + 4C_v C_P} \right),$$

and $H_N^{n+1} = C_P - B_P Q_N^{n+1}$. For $\tau = 0$ the root degenerates to $Q_N^{n+1} = 0$ and $H_N^{n+1} = C_P$.

The initial condition is the steady state at discharge Q_0 with head varying linearly under the Darcy–Weisbach loss. The scheme is exact for the frictionless linear system, and the only discretisation error arises from the first-order treatment of the friction term, which is confirmed to be negligible by the grid independence study of Section 6.3.

APPENDIX C. SUPPORTING DETAILS FOR SECTION 5

C.1 The optimisation over δ . In the proof of Theorem 1 the bound $\mathcal{E}(t)^{1/2} \leq \delta + u + v/\delta$ was obtained with $u = \frac{1}{\sqrt{2}} \int_0^t R^{1/2}$ and $v = \frac{M}{2} \int_0^t b$. The function $\delta \mapsto \delta + v/\delta$ attains its minimum $2\sqrt{v}$ at $\delta = \sqrt{v}$ on $(0, \infty)$, and the bound holds for every $\delta > 0$, so it holds at the minimiser. This yields $\mathcal{E}(t)^{1/2} \leq u + 2\sqrt{v} = u + (2M \int_0^t b)^{1/2}$ as stated. The step is the standard device for converting a differential inequality with an additive forcing term into an integral bound without invoking Grönwall’s inequality, and it is exactly what preserves the $\sqrt{t}$ rather than exponential dependence.

C.2 Monotonicity of the friction nonlinearity. Let $\phi(s) = s|s|$ and $s, \sigma \in \mathbb{R}$. If $s, \sigma \geq 0$ then $\phi(s) - \phi(\sigma) = (s - \sigma)(s + \sigma)$ and the product $(\phi(s) - \phi(\sigma))(s - \sigma) = (s - \sigma)^2(s + \sigma)$ is nonnegative. If $s, \sigma \leq 0$ then $\phi(s) - \phi(\sigma) = -(s - \sigma)(s + \sigma)$ and the product equals $-(s - \sigma)^2(s + \sigma)$, again nonnegative since $s + \sigma \leq 0$. If the signs differ, say $s > 0 > \sigma$, then $\phi(s) > 0 > \phi(\sigma)$ and $s - \sigma > 0$, so the product is positive. Equality forces $s = \sigma$. The map is not Lipschitz on $\mathbb{R}$ but is Lipschitz with constant $2\bar{q}$ on $[-\bar{q}, \bar{q}]$, which is the form used in Corollary 6.

C.3 Derivative bounds for tanh networks. Let $y^{(0)} = z$ and $y^{(l)} = \tanh(W_l y^{(l-1)} + b_l)$. Since $|\tanh^{(k)}| \leq C_k$ on $\mathbb{R}$ for every k , and since the pre-activation is affine in $y^{(l-1)}$, the Faà di Bruno formula expresses a mixed partial of order k of $y^{(l)}$ as a finite sum of terms, each a product of one derivative of $\tanh$ of order at most k with mixed partials of $y^{(l-1)}$ whose orders sum to k , and each differentiation of the pre-activation contributes a factor bounded by $\|W_l\|_\infty$. Induction on l gives $\|\partial^\alpha y^{(L)}\|_\infty \leq C(|\alpha|, L) \prod_{l=1}^L \|W_l\|_\infty^{|\alpha|}$. The Hardy–Krause variation of a C^d function on the unit cube is bounded by the sum over subsets $u \subseteq \{1, \dots, d\}$ of the L^1 norms of $\partial^u F$, hence is finite and controlled by the same product. In practice $\prod_l \|W_l\|_\infty$ is logged during training; for the trained models of Section 7 it settles between 10^2 and 10^3 .

C.4 Koksma–Hlawka. For F of bounded variation in the sense of Hardy and Krause on $[0, 1]^d$ and any point set $\mathcal{Z}$ of size N , $|\int F - \frac{1}{N} \sum_j F(z_j)| \leq V(F) D_N^*(\mathcal{Z})$. For the first N points of a Sobol sequence, $D_N^* \leq C_d N^{-1} (\log N)^{d-1}$ with C_d depending on the generating matrices. Scrambling preserves this bound in expectation and improves the root mean square error to $O(N^{-3/2} (\log N)^{(d-1)/2})$ for smooth integrands, which is why the scrambled variant is used.

APPENDIX D. HYPERPARAMETERS

Setting	Value
Hidden layers	6
Units per layer	128
Activation	tanh
Fourier frequency pairs m	64
Fourier scale σ_F	3.0
Trainable parameters	99,330
Deployed model size, single precision	0.40 MB
Initialisation	Glorot normal
Collocation points N_{col}	4×10^4 , scrambled Sobol, resampled every 10^3 iterations
Valve points N_{bc}	4×10^3
Data points N_d	5×10^3
Phase 1 optimiser	Adam, $\beta = (0.9, 0.999)$
Phase 1 iterations	2×10^4
Phase 1 learning rate	10^{-3} , cosine decay to 10^{-5}
Phase 2 optimiser	L-BFGS, history 50, strong Wolfe line search
Phase 2 iterations	up to 1.5×10^4 , gradient norm tolerance 10^{-8}
Precision	single in Phase 1, double in Phase 2
Weight update interval M	10^3
Weight update rate α	0.1
Causal bins N_s	32
Causal parameter ϵ	scheduled over $\{10^{-2}, 10^{-1}, 1, 10\}$
Data balance ρ	1.0
Parametric family dimension p	4
Training parameter vectors	512 Sobol
Held-out parameter vectors	128 Sobol
Runs per reported result	10
Hardware	one NVIDIA A100, 40 GB