

Initial Value Problem Uncertainty Propagation

Mathias Van Gompel

Leuven Gravity Institute, Department of Electrical Engineering (PSI), KU Leuven, Belgium

Tom Colemont

Leuven Gravity Institute, Department of Electrical Engineering (STADIUS), KU Leuven, Belgium

Tjonnje G.-F. Li

Leuven Gravity Institute, Department of Physics & Astronomy, KU Leuven, Belgium

Johan A.K. Suykens

Department of Electrical Engineering (STADIUS), KU Leuven, Belgium

Frederik De Ceuster

Leuven Gravity Institute, Department of Physics & Astronomy, KU Leuven, Belgium

Abstract

Probabilistic ODE solvers quantify numerical uncertainty by returning a posterior distribution over the solution, rather than only a point estimate. However, these methods typically assume that the initial value problem (IVP) itself is deterministic and fully known. When initial conditions or parameters are uncertain, increasing the numerical accuracy during the solve should not reduce uncertainty about the IVP itself. Standard probabilistic ODE solvers do not distinguish clearly between numerical uncertainty and IVP uncertainty, and may therefore contract the latter inappropriately. Recent work addressed this issue by combining filtering-based probabilistic ODE solvers with numerical quadrature to marginalise correctly over the IVP uncertainty. We extend this work by formulating this outer marginalisation problem in a Bayesian quadrature framework, allowing uncertainty from the quadrature approximation itself to be quantified alongside propagated IVP uncertainty and conditional solver uncertainty. In addition, for Gaussian uncertainty in the initial conditions or parameters, we derive closed-form recursions that can be incorporated directly into filtering and smoothing methods, thereby avoiding numerical quadrature altogether.

Proceedings of the 2nd International Conference on Probabilistic Numerics (ProbNum) 2026, Lappeenranta, Finland. PMLR Volume 341. Copyright 2026 with the author(s)

1 Introduction

Many models in science and engineering describe how a system changes over time. Examples include population growth, mechanical motion, planetary motion, chemical reactions and biological dynamics (Damour and Deruelle, 1985; Daun et al., 2008; Gillespie, 1977; Kermack and McKendrick, 1927). In such models, the current state of the system determines its future evolution through a differential equation.

In practice, this evolution is rarely known in closed form. Even when the model equations are specified exactly, the solution usually has to be approximated numerically. This introduces numerical discretisation uncertainty: the computed solution only approximates the true solution.

Consider the ordinary differential equation (ODE) initial value problem (IVP)

$$\frac{d}{dt}\mathbf{x}(t) = \mathbf{f}(t, \mathbf{x}(t); \boldsymbol{\theta}), \quad t \in [t_0, T], \quad (1)$$

$$\mathbf{x}(t_0) = \mathbf{x}_0. \quad (2)$$

where $\mathbf{x}(t) \in \mathbb{R}^n$ denotes the state trajectory.

Given a vector field $\mathbf{f}$, an initial time t_0 , parameters $\boldsymbol{\theta}$, and an initial condition $\mathbf{x}_0$, we seek a trajectory $\mathbf{x}(t)$ that satisfies the differential equation for all $t \in [t_0, T]$, assuming its existence and uniqueness.

Probabilistic ODE solvers enrich this approximation with uncertainty quantification (Schober et al., 2014, 2019; Tronarp et al., 2019). Rather than providing just a single numerical trajectory, they return a probability distribution over trajectories. By choosing a Gauss–Markov process as a prior over the unknown solution,

one obtains efficient filtering-based formulations in which inference is performed by conditioning on the information that the ODE residual vanishes along the solution trajectory (Schober et al., 2019; Tronarp et al., 2019).

At this point, it is useful to distinguish two quite different roles that uncertainty can play. On the one hand, there is *numerical uncertainty*: even when the inputs $(\mathbf{x}_0, \boldsymbol{\theta})$ for the ODE solver are fixed and known exactly, discretisation and approximation errors remain because the exact flow map is unavailable. On the other hand, there is *IVP uncertainty*: in many applications, neither the initial condition $\mathbf{x}_0$ nor the model parameters $\boldsymbol{\theta}$ are known exactly. If these inputs are uncertain, the uncertainty in the final state must reflect both the numerical uncertainty due to the approximate solver and the IVP uncertainty induced by the uncertain IVP.

This distinction is especially important in probabilistic numerics (PN), which aims to represent uncertainty about numerical computations. (Hennig et al., 2022; Oates and Sullivan, 2019). Standard probabilistic ODE solvers quantify uncertainty conditional on a deterministic, fully specified IVP. If the inputs, i.e. the initial condition or parameters, are uncertain, one faces an additional outer problem: propagating the input distribution through the conditional numerical solver.

It is therefore helpful to separate three conceptually distinct uncertainty sources:

1. **conditional solver uncertainty**, which represents numerical error in the inner ODE solve;
2. **propagated IVP uncertainty**, which reflects uncertainty induced by the uncertain initial conditions or parameters; and

3. **quadrature uncertainty**, which represents uncertainty due to the numerical approximation of the marginalisation over uncertain IVP inputs.

The central principle of this paper can hence be stated as follows:

Uncertainty induced by computation should decrease with numerical refinement, while uncertainty inherited from the IVP should be propagated separately.

Improving the inner numerical solve should reduce conditional solver uncertainty, but not the uncertainty of the IVP itself. This is not a limitation of ODE filtering itself, since an ODE filter is designed to solve a single, fixed IVP. The problem arises only if its conditional covariance is mistakenly interpreted as already containing uncertainty about unknown initial conditions or parameters. Those inputs instead require the explicit outer marginalisation developed in (Yao et al., 2025).

This motivates a distinction between two tasks:

1. **state estimation**: infer hidden states from genuinely informative noisy observations (related to numerical uncertainty);
2. **uncertainty propagation**: push uncertainty in initial conditions or parameters through the dynamics (related to IVP uncertainty).

Yao et al. (2025) showed how uncertain IVP inputs can be handled recursively by combining filtering-based probabilistic ODE solvers with deterministic outer quadrature. For each fixed input, the conditional solver is run and its output is averaged over the input distribution. We make two additions, presented in that order. First, we derive exact closed-form moment recursions for jointly linear–Gaussian models and use their local linearisation as a first-order method for nonlinear IVPs. Second, we formulate the outer marginalisation using Bayesian quadrature (BQ), which yields an explicit outer uncertainty term Σ_k^{out} . The former avoids quadrature and is inexpensive when linearisation is adequate; the latter requires multiple conditional solves but better resolves nonlinear input-to-solution maps.

The main idea can be summarised as follows. For every fixed choice of uncertain inputs $\mathbf{u} = (\mathbf{x}_0, \boldsymbol{\theta})$, we run an ordinary conditional probabilistic ODE solver. This produces a conditional posterior distribution over the solver state. The actual propagation problem then consists of averaging these conditional posteriors over the prior distribution of $\mathbf{u}$.

This point is conceptually important: we do *not* alter the Bayesian inference performed by the solver itself. Rather, we separate two levels: (1) On the inner level, the ODE filter solves a *conditional* inference problem: for fixed inputs $\mathbf{u}$, it computes the conditional distribution $p(\mathbf{x} | \mathbf{u})$. (2) On the outer level, if the inputs are uncertain with distribution $p(\mathbf{u})$, one must additionally propagate this uncertainty through the solver, yielding the marginal distribution

$$p(\mathbf{x}) = \int p(\mathbf{x} | \mathbf{u}) p(\mathbf{u}) d\mathbf{u}. \quad (3)$$

This hierarchical structure is summarised in Fig. 1.

Section 2 reviews filtering-based probabilistic ODE solvers, distinguishes conditional inference from IVP uncertainty propagation and illustrates this distinction through a linear–Gaussian toy model. The common moment decomposition is introduced in Section 3. We then present the first-order method in Section 4, followed by BQ in Section 5.

2 Background

2.1 Filtering-based probabilistic ODE solvers

Let $t_0 < t_1 < \dots < t_K$ be a discretisation grid. A filtering-based probabilistic ODE solver introduces a solver state $\mathbf{x}_k$ and a solver-internal residual variable $\mathbf{y}_k$ such that

$$\mathbf{x}_k | \mathbf{x}_{k-1} \sim p(\mathbf{x}_k | \mathbf{x}_{k-1}), \quad (4)$$

$$\mathbf{y}_k | \mathbf{x}_k \sim p(\mathbf{y}_k | \mathbf{x}_k). \quad (5)$$

Equation (4) describes the Gauss–Markov prior with Markovian kernel, while (5) introduces the solver-internal residual or consistency condition.

For ODE filters, the state $\mathbf{x}_k$ typically contains the solution and several of its derivatives at time t_k , while $\mathbf{y}_k$ is often a residual that compares a predicted derivative against a vector-field evaluation and is targeted at zero.

Because $\mathbf{y}_k$ is solver-internal rather than external data, conditioning on it should reduce only numerical uncertainty for fixed IVP inputs. It must not reduce uncertainty arising from genuinely unknown IVP inputs such as $\mathbf{x}_0$ or $\boldsymbol{\theta}$. For a more detailed introduction of probabilistic ODE solvers, we refer to (Schober et al., 2019; Tronarp et al., 2019).

2.2 Conditional inference versus uncertainty propagation

Let

$$\mathbf{u} := \begin{bmatrix} \mathbf{x}_0 \\ \boldsymbol{\theta} \end{bmatrix}$$

collect the uncertain IVP inputs, with prior distribution $p(\mathbf{u})$. For every fixed value of $\mathbf{u}$, the probabilistic solver returns a conditional posterior

$$p(\mathbf{x}_k | \mathbf{y}_{1:k}, \mathbf{u}). \quad (6)$$

In most practical filters, this conditional posterior is approximated as a Gaussian distribution,

$$\begin{aligned} \mathbf{x}_k | \mathbf{y}_{1:k}, \mathbf{u} &\sim p(\mathbf{x}_k | \mathbf{y}_{1:k}, \mathbf{u}) \\ &\approx \mathcal{N}(\mathbf{m}_k(\mathbf{u}), \bar{\mathbf{P}}_k(\mathbf{u})), \end{aligned} \quad (7)$$

so that the conditional solver output is described by a mean map $\mathbf{m}_k(\mathbf{u})$ and a covariance map $\bar{\mathbf{P}}_k(\mathbf{u})$.

The physically meaningful propagated distribution is then obtained by averaging these conditional posteriors over the uncertain input distribution:

$$p_{\text{prop}}(\mathbf{x}_k | \mathbf{y}_{1:k}) := \int p(\mathbf{x}_k | \mathbf{y}_{1:k}, \mathbf{u}) p(\mathbf{u}) d\mathbf{u}. \quad (8)$$

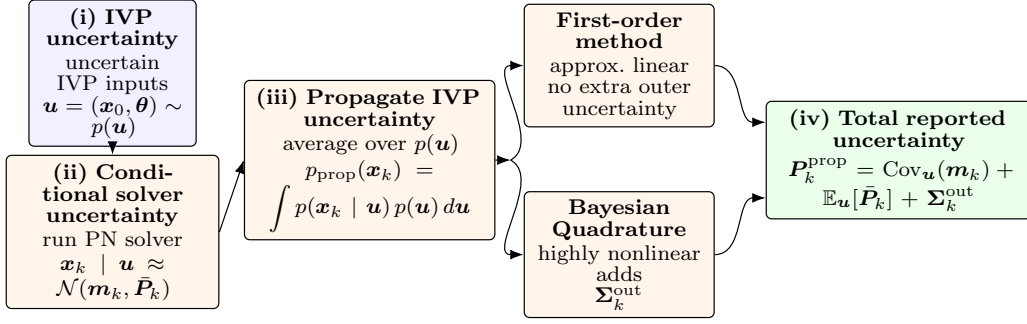


Figure 1: Hierarchical view of uncertainty propagation for IVPs with uncertain inputs. For fixed inputs $u = (x_0, \theta)$, the probabilistic ODE solver returns a conditional distribution with conditional solver uncertainty. IVP uncertainty is then propagated by averaging these conditional outputs over the input distribution $p(u)$. In approximately linear regimes, this outer propagation can be handled by Jacobian recursion without introducing additional outer numerical uncertainty. In strongly nonlinear regimes, the outer marginalisation must generally be approximated numerically, for example by quadrature, Monte Carlo, or Bayesian quadrature. BQ additionally reports an outer computational uncertainty term.

This equation is one of the key objects in the paper. It simply states that the uncertainty at time t_k should account for *all plausible inputs*, weighted by the probability assigned to them initially. Put differently, the final propagated distribution is an outer marginalisation of conditional solver outputs.

The integral in Eq. (8) is generally intractable. Even if each conditional posterior is Gaussian, the marginal need not admit a simple analytic form, because the maps $u \mapsto m_k(u)$ and $u \mapsto \bar{P}_k(u)$ may depend nonlinearly on the uncertain inputs. This is exactly the point at which approximation methods such as Monte Carlo, deterministic quadrature, or Bayesian quadrature become necessary (O’Hagan, 1991; Briol et al., 2019).

2.3 Toy model

In this section, we use a minimal linear–Gaussian toy model to illustrate why ordinary filtering uncertainty does not, in general, provide the correct uncertainty propagation for an IVP with uncertain initial conditions. The example isolates the distinction between solver-induced uncertainty, which should contract under numerical conditioning, and uncertainty inherited from the initial condition, which should remain present in the propagated solution.

Consider a single time step with the scalar linear–Gaussian model

$$\begin{aligned} x_0 &\sim \mathcal{N}(0, \sigma_0^2), & x_1 &= x_0 + w, & w &\sim \mathcal{N}(0, q), \\ y_1 &= x_1 + v, & v &\sim \mathcal{N}(0, r), \end{aligned} \quad (9)$$

fixed solver-internal residual value $y_1 \in \mathbb{R}$.

If we interpret this as an ordinary filtering problem, then the Kalman update yields posterior variance (Särkkä and Svensson, 2023)

$$P_{\text{filter}} = \frac{r(\sigma_0^2 + q)}{\sigma_0^2 + q + r}. \quad (10)$$

This is the familiar variance contraction from Bayesian conditioning:

$$\sigma_0^2 \rightarrow \infty \quad \Rightarrow \quad P_{\text{filter}} \rightarrow r. \quad (11)$$

This limit is inappropriate for IVP uncertainty propagation. Here y_1 is not informative external data about the unknown initial condition; it is a solver-internal consistency condition enforcing that the numerical step approximately satisfies the ODE. Conditioning on y_1 should therefore reduce only numerical uncertainty and must not eliminate genuine uncertainty about the IVP input x_0 . Accordingly, the correct propagated behaviour is

$$\sigma_0^2 \rightarrow \infty \quad \Rightarrow \quad P_{\text{prop}} \rightarrow \infty. \quad (12)$$

The toy model therefore exposes the core tension clearly: ordinary filtering treats y_1 as informative about the latent state itself, whereas IVP uncertainty propagation should treat it only as a numerical consistency condition. We now formalise this propagation principle recursively.

Propagation prediction distribution Assume that at time $k - 1$ we have already propagated both numerical uncertainty and IVP uncertainty via $p_{\text{prop}}(x_{k-1} | y_{1:k-1})$. The prediction step is the standard Chapman–Kolmogorov update (Särkkä and Svensson, 2023; Yao et al., 2025).

$$p_{\text{prop}}(x_k | y_{1:k-1}) = \quad (13)$$

$$\int p(x_k | x_{k-1}) p_{\text{prop}}(x_{k-1} | y_{1:k-1}) dx_{k-1}, \quad (14)$$

Propagation filtering distribution The update step is where the distinction from ordinary filtering enters. For a *fixed* previous state x_{k-1} , Bayes’ rule yields the local conditional update

$$p(x_k | y_k, x_{k-1}) = \frac{p(y_k | x_k) p(x_k | x_{k-1})}{p(y_k | x_{k-1})}, \quad (15)$$

where

$$p(y_k | x_{k-1}) = \int p(y_k | x_k) p(x_k | x_{k-1}) dx_k. \quad (16)$$

So far this is still ordinary Bayesian conditioning, but crucially only at the local level for fixed x_{k-1} . The propagation update is then obtained by averaging these locally updated conditionals over the previously propagated uncertainty:

$$\begin{aligned}
p_{\text{prop}}(\mathbf{x}_k \mid \mathbf{y}_{1:k}) &= \\
&\int p(\mathbf{x}_k \mid \mathbf{y}_k, \mathbf{x}_{k-1}) p_{\text{prop}}(\mathbf{x}_{k-1} \mid \mathbf{y}_{1:k-1}) d\mathbf{x}_{k-1} = \\
&\int \frac{p(\mathbf{y}_k \mid \mathbf{x}_k) p(\mathbf{x}_k \mid \mathbf{x}_{k-1})}{p(\mathbf{y}_k \mid \mathbf{x}_{k-1})} p_{\text{prop}}(\mathbf{x}_{k-1} \mid \mathbf{y}_{1:k-1}) d\mathbf{x}_{k-1}.
\end{aligned}
\tag{17}$$

$$\tag{18}$$

Remark 1. For comparison, the standard Bayesian filter associated with the state-space model normalises globally. In contrast, the propagation update (18) normalises locally with $p(\mathbf{y}_k \mid \mathbf{x}_{k-1})$ before averaging over $\mathbf{x}_{k-1}$.

3 Propagated moments

The conditional solver is available for every fixed input $\mathbf{u}$, but the outer average in Eq. (8) is generally intractable. Both methods developed below in Section 4 and Section 5 target the same first two moments, so we state those quantities before introducing either approximation.

For this reason, it is useful to separate the computational task into two layers:

1. for fixed $\mathbf{u}$, apply the conditional probabilistic ODE solver and obtain $(\mathbf{m}_k(\mathbf{u}), \bar{\mathbf{P}}_k(\mathbf{u}))$;
2. evaluate or approximate the outer integrals over $p(\mathbf{u})$.

This modular decomposition leaves the inner solver unchanged. The outer layer may be treated analytically, by first-order linearisation, by Monte Carlo, by deterministic quadrature as in (Yao et al., 2025), or by BQ.

3.1 Gaussian-weighted integrals

Assume a Gaussian distribution

$$\mathbf{u} := \begin{bmatrix} \mathbf{x}_0 \\ \boldsymbol{\theta} \end{bmatrix} \in \mathbb{R}^{d_u}, \quad \mathbf{u} \sim \mathcal{N}(\mathbf{m}_u, \mathbf{P}_u). \tag{19}$$

For each deterministic input $\mathbf{u}$, apply the conditional solver to obtain

$$\mathbf{x}_k \mid \mathbf{u} \approx \mathcal{N}(\mathbf{m}_k(\mathbf{u}), \bar{\mathbf{P}}_k(\mathbf{u})), \tag{20}$$

where $\mathbf{m}_k(\mathbf{u})$ is the conditional solver mean and $\bar{\mathbf{P}}_k(\mathbf{u})$ is the corresponding conditional covariance.

The first two moments of the propagation distribution then follow from the law of total expectation and the law of total variance:

$$\mathbf{m}_k^{\text{prop}} = \mathbb{E}_{\mathbf{u}}[\mathbf{m}_k(\mathbf{u})], \tag{21}$$

$$\mathbf{P}_k^{\text{prop}} = \mathbb{E}_{\mathbf{u}}[\bar{\mathbf{P}}_k(\mathbf{u})] + \text{Cov}_{\mathbf{u}}(\mathbf{m}_k(\mathbf{u})). \tag{22}$$

These two formulas separate the two contributions to the propagated covariance. The first term is the average

conditional solver covariance and reflects conditional numerical uncertainty. The second term measures how strongly the conditional mean varies across different plausible inputs and therefore captures the spread induced by uncertain initial conditions and parameters.

If $\mathbf{m}_k(\mathbf{u})$ hardly changes with $\mathbf{u}$, then propagated IVP uncertainty is small. If it varies strongly with $\mathbf{u}$, then propagated IVP uncertainty becomes large even if each conditional solver covariance is itself small.

Remark 2 (Consistency under solver refinement). Suppose the numerical conditional mean and covariance converge pointwise to those induced by the exact flow as the inner solver is refined and suppose they are dominated by integrable bounds under $p(\mathbf{u})$. Dominated convergence then implies convergence of both moments in Eqs. (21) and (22) to the moments of the exact pushforward. Refinement should therefore improve the approximation of the propagated distribution without collapsing its intrinsic IVP uncertainty.

4 First-order propagation

We first treat the setting in which outer marginalisation is exact. The same algebra then yields a local first-order approximation for nonlinear IVPs, making explicit both its computational appeal and its limitations.

4.1 Exact jointly linear–Gaussian recursion

Assume a linear–Gaussian model

$$\boldsymbol{\theta} \sim \mathcal{N}(\mathbf{m}_{\boldsymbol{\theta}}, \mathbf{P}_{\boldsymbol{\theta}}), \tag{23}$$

$$\mathbf{x}_0 \sim \mathcal{N}(\mathbf{m}_0, \mathbf{P}_0), \tag{24}$$

$$\mathbf{x}_k = \mathbf{A}_{k-1}\mathbf{x}_{k-1} + \mathbf{B}_{k-1}\boldsymbol{\theta} + \mathbf{q}_{k-1}, \tag{25}$$

$$\mathbf{y}_k = \mathbf{C}_k\mathbf{x}_k + \mathbf{D}_k\boldsymbol{\theta} + \mathbf{r}_k, \tag{26}$$

where

$$\mathbf{q}_{k-1} \sim \mathcal{N}(\mathbf{0}, \mathbf{Q}_{k-1}), \quad \mathbf{r}_k \sim \mathcal{N}(\mathbf{0}, \mathbf{R}_k). \tag{27}$$

Write the joint prior of $(\mathbf{x}_0, \boldsymbol{\theta})$ as

$$\begin{bmatrix} \mathbf{x}_0 \\ \boldsymbol{\theta} \end{bmatrix} \sim \mathcal{N}\left(\begin{bmatrix} \mathbf{m}_0 \\ \mathbf{m}_{\boldsymbol{\theta}} \end{bmatrix}, \Sigma_0\right), \quad \Sigma_0 := \begin{bmatrix} \mathbf{P}_0 & \mathbf{0} \\ \mathbf{0} & \mathbf{P}_{\boldsymbol{\theta}} \end{bmatrix}. \tag{28}$$

So we assume that $\mathbf{x}_0$ and $\boldsymbol{\theta}$ are independent.

For fixed $(\mathbf{x}_0, \boldsymbol{\theta})$, the standard Kalman recursions are

$$\mathbf{m}_0(\mathbf{x}_0, \boldsymbol{\theta}) = \mathbf{x}_0, \tag{29}$$

$$\bar{\mathbf{P}}_0 = \mathbf{0}, \tag{30}$$

$$\mathbf{m}_k^-(\mathbf{x}_0, \boldsymbol{\theta}) = \mathbf{A}_{k-1}\mathbf{m}_{k-1}(\mathbf{x}_0, \boldsymbol{\theta}) + \mathbf{B}_{k-1}\boldsymbol{\theta}, \tag{31}$$

$$\bar{\mathbf{P}}_k^- = \mathbf{A}_{k-1}\bar{\mathbf{P}}_{k-1}\mathbf{A}_{k-1}^\top + \mathbf{Q}_{k-1}, \tag{32}$$

$$\mathbf{v}_k(\mathbf{x}_0, \boldsymbol{\theta}) = \mathbf{y}_k - \mathbf{C}_k\mathbf{m}_k^-(\mathbf{x}_0, \boldsymbol{\theta}) - \mathbf{D}_k\boldsymbol{\theta}, \tag{33}$$

$$\mathbf{S}_k = \mathbf{C}_k\bar{\mathbf{P}}_k^-\mathbf{C}_k^\top + \mathbf{R}_k, \tag{34}$$

$$\mathbf{K}_k = \bar{\mathbf{P}}_k^-\mathbf{C}_k^\top\mathbf{S}_k^{-1}, \tag{35}$$

$$\mathbf{m}_k(\mathbf{x}_0, \boldsymbol{\theta}) = \mathbf{m}_k^-(\mathbf{x}_0, \boldsymbol{\theta}) + \mathbf{K}_k\mathbf{v}_k(\mathbf{x}_0, \boldsymbol{\theta}), \tag{36}$$

$$\bar{\mathbf{P}}_k = \bar{\mathbf{P}}_k^- - \mathbf{K}_k\mathbf{S}_k\mathbf{K}_k^\top. \tag{37}$$

such that

$$\mathbf{x}_k \mid \mathbf{y}_{1:k}, \mathbf{x}_0, \boldsymbol{\theta} \sim \mathcal{N}(\mathbf{m}_k(\mathbf{x}_0, \boldsymbol{\theta}), \bar{\mathbf{P}}_k), \quad (38)$$

where $\bar{\mathbf{P}}_k$ and the Kalman gain $\mathbf{K}_k$ do not depend on the values of $(\mathbf{x}_0, \boldsymbol{\theta})$. Consequently, the conditional mean is affine in the uncertain inputs and the outer marginalisation remains Gaussian.

4.2 Closed-form propagated moments

Let $\mathbf{m}_k^*$ be the conditional mean obtained at the prior mean input. Then

$$\mathbf{m}_k(\mathbf{x}_0, \boldsymbol{\theta}) = \mathbf{m}_k^* + \mathbf{J}_k^x(\mathbf{x}_0 - \mathbf{m}_0) + \mathbf{J}_k^\theta(\boldsymbol{\theta} - \mathbf{m}_\theta). \quad (39)$$

These Jacobians satisfy the recursions

$$\mathbf{J}_0^x = \mathbf{I}_d, \quad (40)$$

$$\mathbf{J}_k^{x,-} = \mathbf{A}_{k-1} \mathbf{J}_{k-1}^x, \quad (41)$$

$$\mathbf{J}_k^x = (\mathbf{I}_d - \mathbf{K}_k \mathbf{C}_k) \mathbf{J}_k^{x,-}, \quad (42)$$

$$\mathbf{J}_0^\theta = 0, \quad (43)$$

$$\mathbf{J}_k^{\theta,-} = \mathbf{A}_{k-1} \mathbf{J}_{k-1}^\theta + \mathbf{B}_{k-1}, \quad (44)$$

$$\mathbf{J}_k^\theta = (\mathbf{I}_d - \mathbf{K}_k \mathbf{C}_k) \mathbf{J}_k^{\theta,-} - \mathbf{K}_k \mathbf{D}_k. \quad (45)$$

The two Jacobian blocks have distinct meanings. The matrix $\mathbf{J}_k^x$ transports uncertainty already present in the initial state, whereas $\mathbf{J}_k^\theta$ measures how parameter perturbations affect the state throughout the dynamics. Initial-condition and parameter uncertainty therefore share an outer marginalisation but are not assumed to propagate identically. Since the model is linear, this first-order representation is exact rather than approximate. More precisely, it is exact for the linear–Gaussian state-space model to which the Kalman recursions are applied. When this model is obtained by locally linearising a nonlinear ODE filter, the propagated moments inherit the approximation error of that linearisation. Moreover, the Gaussian closure represents uncertainty through a covariance ellipsoid that is symmetric about the mean. A nonlinear input-to-solution map may instead produce asymmetric or otherwise non-Gaussian pushforward distributions, so a first-order Gaussian approximation can underestimate or misorient the propagated uncertainty.

As a consequence,

$$\mathbf{m}_k^{\text{prop}} = \mathbf{m}_k^*, \quad (46)$$

and if we write $\mathbf{J}_k := [\mathbf{J}_k^x \quad \mathbf{J}_k^\theta]$, then

$$\mathbf{P}_k^{\text{prop}} = \bar{\mathbf{P}}_k + \mathbf{J}_k \boldsymbol{\Sigma}_0 \mathbf{J}_k^\top. \quad (47)$$

This is the exact outer marginalisation for the approximating linear–Gaussian solver model; no quadrature uncertainty is introduced.

4.3 Extension to nonlinear IVPs

For a nonlinear IVP, the preceding closed-form result is not exact. We instead linearise the conditional mean

map around the input mean $\mathbf{m}_u$:

$$\begin{aligned} \mathbf{m}_k(\mathbf{u}) &\approx \mathbf{m}_k(\mathbf{m}_u) + \mathbf{J}_k(\mathbf{u} - \mathbf{m}_u), \\ \mathbf{J}_k &:= \left. \frac{\partial \mathbf{m}_k(\mathbf{u})}{\partial \mathbf{u}} \right|_{\mathbf{u}=\mathbf{m}_u}. \end{aligned} \quad (48)$$

The Jacobian is propagated through the locally linearised prediction and update maps of the conditional ODE filter, in direct analogy with an extended Kalman filter (Tronarp et al., 2019). The resulting moment approximation is

$$\mathbf{m}_k^{\text{FO}} = \mathbf{m}_k(\mathbf{m}_u), \quad \mathbf{P}_k^{\text{FO}} = \bar{\mathbf{P}}_k(\mathbf{m}_u) + \mathbf{J}_k \mathbf{P}_u \mathbf{J}_k^\top. \quad (49)$$

Thus the nonlinear method requires one conditional solve at the mean input and a sensitivity recursion. It is best viewed as an extended-Kalman-style approximation: it is reliable for smooth vector fields, small steps and concentrated input distributions whose input-to-solution map is nearly affine. Curvature and state interactions can make it underestimate or misorient the propagated covariance, motivating BQ.

5 Bayesian quadrature

When the first-order approximation is inadequate, the outer integrals can be evaluated at several input locations. BQ treats the conditional mean map as an unknown function under a Gaussian-process prior and turns its evaluations into a posterior distribution over the required Gaussian-weighted moments.

5.1 Bayesian quadrature moment transform

The goal is to approximate the outer Gaussian-weighted integrals defining the propagated moments while assigning an additional numerical-uncertainty term under the BQ model. Monte Carlo methods target the same moments by drawing $\mathbf{u}^{(i)} \sim p(\mathbf{u})$, running a solver for each sample and replacing the expectations in Eqs. (21) and (22) by empirical averages. The BQ construction used here instead combines deterministic sigma points with Gaussian-process quadrature weights and returns an additional uncertainty term for the outer numerical integration (Prüher and Šimandl, 2016).

A detailed derivation appears in (Prüher and Šimandl, 2016). Write the uncertain input as

$$\begin{aligned} \mathbf{u} &= \mathbf{m}_u + \mathbf{L}_u \boldsymbol{\xi}, \quad \boldsymbol{\xi} \sim \mathcal{N}(0, \mathbf{I}_{d_u}), \\ \mathbf{L}_u \mathbf{L}_u^\top &= \mathbf{P}_u. \end{aligned} \quad (50)$$

Define the propagated conditional-mean map

$$g_k(\boldsymbol{\xi}) := \mathbf{m}_k(\mathbf{m}_u + \mathbf{L}_u \boldsymbol{\xi}). \quad (51)$$

The required quantities are $\mathbf{m}_k^{\text{prop}} = \mathbb{E}_{\boldsymbol{\xi}}[g_k(\boldsymbol{\xi})]$ and $\boldsymbol{\Pi}_k = \text{Cov}_{\boldsymbol{\xi}}(g_k(\boldsymbol{\xi}))$, the propagated IVP covariance.

Choose sigma points $\{\boldsymbol{\xi}^{(i)}\}_{i=1}^N$, set $\mathbf{u}^{(i)} = \mathbf{m}_u + \mathbf{L}_u \boldsymbol{\xi}^{(i)}$ and run the conditional solver at each input. We use tensor-product Gauss–Hermite points where, in classical quadrature, we would use normalised deterministic

weights $\bar{w}_i$ (Stroud et al., 1966). Collect the conditional means row-wise in

$$\mathbf{G}_k \in \mathbb{R}^{N \times d_x}, \quad (\mathbf{G}_k)_{i,:} = g_k(\boldsymbol{\xi}^{(i)})^\top. \quad (52)$$

Each output component of g_k is modelled by an independent GP with shared kernel k (Prüher and Šimandl, 2016; Särkkä et al., 2015). Its Gram matrix is

$$\mathbf{K}_{ij} := k(\boldsymbol{\xi}^{(i)}, \boldsymbol{\xi}^{(j)}). \quad (53)$$

Define the kernel mean and second-moment matrices by

$$\boldsymbol{\ell}_i := \mathbb{E}_{\boldsymbol{\xi}}[k(\boldsymbol{\xi}, \boldsymbol{\xi}^{(i)})]. \quad (54)$$

$$\mathbf{L}_{ij} := \mathbb{E}_{\boldsymbol{\xi}}[k(\boldsymbol{\xi}, \boldsymbol{\xi}^{(i)})k(\boldsymbol{\xi}, \boldsymbol{\xi}^{(j)})]. \quad (55)$$

The BQ weights and moment approximations are

$$\mathbf{w} := \mathbf{K}^{-1}\boldsymbol{\ell}, \quad (56)$$

$$\mathbf{W} := \mathbf{K}^{-1}\mathbf{L}\mathbf{K}^{-1}, \quad (57)$$

$$\boldsymbol{\mu}_k^{\text{BQ}} := \mathbf{G}_k^\top \mathbf{w}, \quad (58)$$

$$\hat{\boldsymbol{\Pi}}_k^{\text{BQ}} := \mathbf{G}_k^\top \mathbf{W} \mathbf{G}_k - \boldsymbol{\mu}_k^{\text{BQ}}(\boldsymbol{\mu}_k^{\text{BQ}})^\top. \quad (59)$$

The subtraction centres the approximated second moment and therefore separates propagated IVP covariance from the BQ integration term. A more detailed derivation can be found in Section B.

With $\bar{k} := \mathbb{E}_{\boldsymbol{\xi}}[k(\boldsymbol{\xi}, \boldsymbol{\xi})]$, the BQ model assigns the integrated posterior variance (Prüher and Šimandl, 2016)

$$\sigma_{\text{int}}^2 := \bar{k} - \text{tr}(\mathbf{K}^{-1}\mathbf{L}). \quad (60)$$

The distinct outer quadrature-uncertainty term is

$$\Sigma_k^{\text{out}} := \sigma_{\text{int}}^2 \mathbf{I}_{d_x}. \quad (61)$$

The average conditional solver covariance uses the BQ weights on the same nodes,

$$\mathbb{E}_{\mathbf{u}}[\widehat{\bar{\mathbf{P}}}_k(\mathbf{u})] := \sum_{i=1}^N w_i \bar{\mathbf{P}}_k(\mathbf{u}^{(i)}). \quad (62)$$

and $\mathbf{m}_k^{\text{prop,BQ}} := \boldsymbol{\mu}_k^{\text{BQ}}$. The covariance decomposes into three interpretable terms:

$$\begin{aligned} \mathbf{P}_k^{\text{prop,BQ}} = & \underbrace{\sum_{i=1}^N w_i \bar{\mathbf{P}}_k(\mathbf{u}^{(i)})}_{\text{conditional solver}} \\ & + \underbrace{\hat{\boldsymbol{\Pi}}_k^{\text{BQ}}}_{\text{propagated IVP}} + \underbrace{\Sigma_k^{\text{out}}}_{\text{outer BQ}}. \end{aligned} \quad (63)$$

For the exponentiated-quadratic kernel, all kernel integrals above are available in closed form. A deterministic solver removes the first term; deterministic Gauss–Hermite quadrature uses $\bar{w}_i$ and has $\Sigma_k^{\text{out}} = 0$; exact linear marginalisation needs neither quadrature term. Tensor-product Gauss–Hermite uses $N = n^{d_u}$ nodes, so both quadrature methods are practical mainly for low-dimensional uncertain inputs.

The distinction between BQ and classical Gauss–Hermite quadrature is shown in Section D. In that comparison,

both methods use the same sigma points and the same conditional solver evaluations, the difference comes from the outer moment weights. Classical Gauss–Hermite uses the deterministic weights $\bar{w}_i$, whereas BQ uses the kernel-dependent weights $\mathbf{w} = \mathbf{K}^{-1}\boldsymbol{\ell}$ for the mean and $\mathbf{W} = \mathbf{K}^{-1}\mathbf{L}\mathbf{K}^{-1}$ for the second moment. Consequently, BQ also returns the additional outer uncertainty term Σ_k^{out} , which is absent in deterministic quadrature.

6 Numerical experiments

This section benchmarks the *first-order* propagation method and BQ against Monte Carlo (MC) on three IVPs. We evaluate the accuracy–runtime tradeoff for the pushforward moments using a high-sample MC reference. The code for the numerical experiments is available online¹.

6.1 Protocol and implementation details

Choice of input distributions and BQ prior. The distributions over $\mathbf{x}_0$ and $\boldsymbol{\theta}$ are modelling assumptions external to the solver and should encode domain knowledge, measurement uncertainty, or an inferred posterior. We use Gaussian distributions for tractability. BQ uses an exponentiated-quadratic kernel because its Gaussian kernel integrals are available in closed form; other kernels are possible when the corresponding integrals can be evaluated analytically.

Benchmark problems and input distributions.

We consider the three benchmark IVPs summarised in Table 1. They span linear and nonlinear dynamics and include uncertainty in both initial conditions and model parameters.

As shown in Table 1, the linear and Lotka–Volterra benchmarks contain uncertain initial conditions, whereas the logistic benchmark contains uncertainty in the carrying-capacity parameter b . These problems exercise qualitatively different propagation mechanisms. For the linear equation, initial-state variance expands with the flow. In the logistic equation, uncertainty in the carrying capacity changes the attracting level rather than merely perturbing the initial state. Lotka–Volterra combines a two-dimensional uncertain initial condition with strongly nonlinear coupled dynamics. The common outer formulation does not assume that these sources propagate identically; their effects enter through the corresponding input-to-solution sensitivities.

Conditional probabilistic ODE solver. For each deterministic input realisation $\mathbf{u}$, we run the same conditional solver: an EKF ODE filter with an integrated Wiener process prior of order $q = 4$ (IWP(4)) and adaptive step-size selection (Schober et al., 2019; Tronarp et al., 2019; Bosch et al., 2021). This yields the conditional mean $\mathbf{m}_k(\mathbf{u})$ and covariance $\bar{\mathbf{P}}_k(\mathbf{u})$ at discrete time points. The first-order method uses a single such

¹ProbNum26.UncertaintyPropagation

Table 1: Benchmark IVPs and uncertain-input distributions used in the numerical experiments.

Problem	Vector field	Time interval	Initial value	Parameters
Linear	$f(t, x) = x$	$t \in [0, 3]$	$x(0) \sim \mathcal{N}(1, 10^{-2})$	—
Logistic	$f(t, x) = ax \left(1 - \frac{x}{b}\right)$	$t \in [0, 3]$	$x(0) = 0.05$	$a = 3, \quad b \sim \mathcal{N}(3, 10^{-2})$
Lotka–Volterra	$\mathbf{f}(t, \mathbf{x}) = \begin{bmatrix} ax_1 - bx_1x_2 \\ -cx_2 + dx_1x_2 \end{bmatrix}$	$t \in [0, 2]$	$\mathbf{x}(0) \sim \mathcal{N}\left(\begin{bmatrix} 5 \\ 5 \end{bmatrix}, 0.3 \mathbf{I}_2\right)$	$(a, b, c, d) = (5, 0.5, 5, 0.5)$

solver run at the nominal input together with Jacobian propagation, whereas BQ evaluates the conditional solver at each sigma point.

Reference solution (high-sample MC). As reference, we use a Monte Carlo estimate with $N_{\text{ref}} = 10^6$: sample $\mathbf{u} \sim p(\mathbf{u})$, solve the deterministic IVP with `solve_ivp` using DOP853 and compute the empirical mean and covariance ($\mathbf{m}_k^{\text{ref}}, \mathbf{P}_k^{\text{ref}}$) at each reported time point.

Methods, baselines and runtime measurement. We compare the first-order method, BQ and MC. The first-order moments use Eq. (49); BQ uses Section 5.1 and Eqs. (59) and (63) and the GH nodes in Table 2. For MC we use

$$N \in \{200, 500, 10^3, 2 \cdot 10^3, 5 \cdot 10^3, 10^4, 2 \cdot 10^4, 5 \cdot 10^4, 10^5\}.$$

For each N , MC uses N deterministic forward solves and empirical moment estimates. Each experiment is repeated 10 times, except for $N = 10^5$, which is run once. Errors are shown as boxplots. Runtime is measured as wall-clock time for the full method evaluation, including all deterministic forward solves for MC and all conditional solver calls for BQ.

Benchmark metric. At each time point t_k , a method returns a Gaussian approximation $x(t_k) \approx \mathcal{N}(\mathbf{m}_k, \mathbf{P}_k)$, which we compare to the Gaussian reference $\mathcal{N}(\mathbf{m}_k^{\text{ref}}, \mathbf{P}_k^{\text{ref}})$ using the closed-form 2-Wasserstein distance (Gelbrich, 1990):

$$W_2^2(\mathcal{N}(\mathbf{m}_1, \mathbf{C}_1), \mathcal{N}(\mathbf{m}_2, \mathbf{C}_2)) = \|\mathbf{m}_1 - \mathbf{m}_2\|^2 + \text{tr}\left(\mathbf{C}_1 + \mathbf{C}_2 - 2(\mathbf{C}_2^{1/2} \mathbf{C}_1 \mathbf{C}_2^{1/2})^{1/2}\right). \quad (64)$$

Define $w_k := W_2(\mathcal{N}(\mathbf{m}_k, \mathbf{P}_k), \mathcal{N}(\mathbf{m}_k^{\text{ref}}, \mathbf{P}_k^{\text{ref}}))$, and aggregate over time by

$$W_{2\text{rms}} := \left(\frac{1}{T} \sum_{k=1}^T w_k^2\right)^{1/2}. \quad (65)$$

Remark 3 (IVP-propagation accuracy). *The runtime-accuracy figures evaluate recovery of the pushforward mean and covariance induced by uncertain inputs, namely $\mathbb{E}_{\mathbf{u}}[\mathbf{m}_k(\mathbf{u})]$ and $\text{Cov}_{\mathbf{u}}(\mathbf{m}_k(\mathbf{u}))$ in Eqs. (21) and (22). Conditional solver covariance and Σ_k^{out} have different meanings and are not the target of this metric.*

6.2 Results

Fig. 2 shows the expected MC convergence trend: error decreases steadily with increasing N at increasing runtime. In this linear setting, the first-order method achieves $W_{2\text{rms}} \approx 10^{-3}$ at sub-second runtime, outperforming MC even at $N = 10^5$ (which attains errors on the order of 10^{-3} – 10^{-2} at two orders of magnitude larger runtime). BQ reaches a comparable error level but at higher cost because it performs multiple conditional solver evaluations. This is consistent with the linear theory: the first-order outer marginalisation is exact for the approximating linear–Gaussian model, while the conditional ODE solve remains numerical.

Logistic growth. In Fig. 3, the first-order method remains highly competitive: it matches or exceeds the accuracy of large- N MC at a fraction of the runtime, which is consistent with the mapping $\mathbf{u} \mapsto \mathbf{m}_k(\mathbf{u})$ being close to affine over the support of the input distribution. BQ further reduces error, down to the 10^{-4} range in these experiments, at a noticeable increase in runtime.

Lotka–Volterra. Fig. 4 highlights the regime where local linearisation becomes unreliable. The first-order method incurs a substantially larger error than MC, consistent with strong nonlinearity and state–state interactions amplifying higher-order effects of IVP uncertainty. BQ, which evaluates the nonlinear conditional mean at multiple input locations, recovers accuracy and improves substantially over MC at comparable or lower runtime

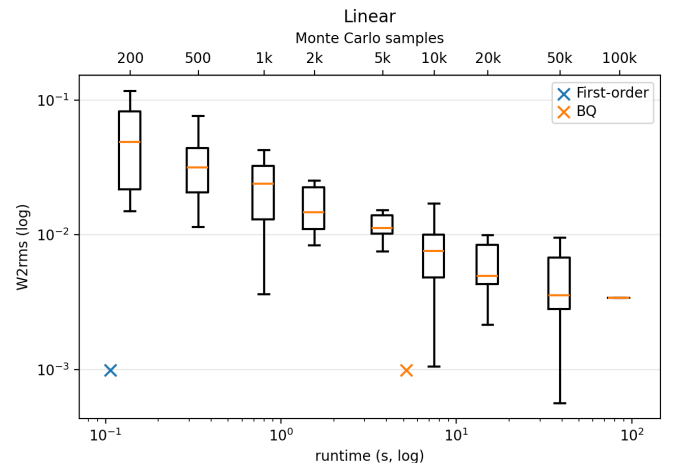


Figure 2: Linear IVP: distribution of MC errors (boxplots) versus runtime, measured by $W_{2\text{rms}}$ relative to a $N_{\text{ref}} = 10^6$ reference.

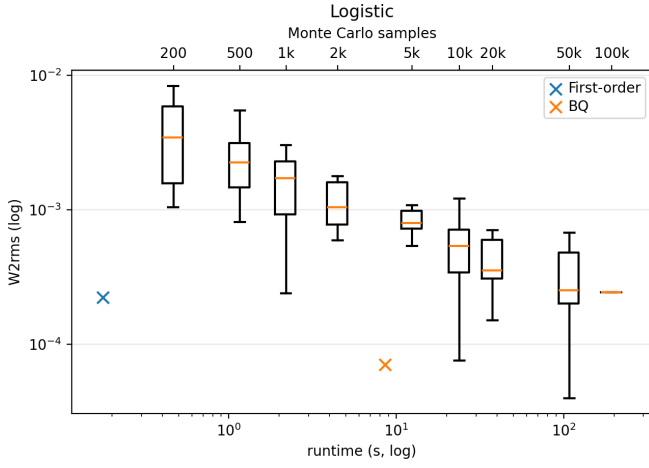


Figure 3: Logistic IVP: distribution of MC errors (boxplots) versus runtime, measured by $W2_{rms}$ relative to a $N_{ref} = 10^6$ reference.

than very large MC sample sizes.

The trajectories of the solved benchmarks and the propagation of IVP uncertainty are given in [Section E](#)

7 Discussion

The first-order method is exact for jointly linear–Gaussian models. In that setting, the conditional mean is affine in the uncertain inputs and the conditional covariance does not depend on their realised values, so the outer marginalisation can be evaluated analytically. For nonlinear problems, it becomes an extended-Kalman-style approximation in which the input-to-solution map is linearised around the mean input. It is therefore most suitable when the input distribution is concentrated and the conditional mean map is approximately affine on its high-probability region.

The experiments illustrate this distinction. First-order propagation is accurate and computationally efficient for the linear and logistic problems, but performs worse for

Lotka–Volterra, where nonlinear state interactions introduce higher-order dependence on the uncertain initial condition. In such cases, the method may underestimate or misorient the propagated covariance. Moreover, the resulting Gaussian approximation only describes the first two moments and may not capture skewness, heavy tails, or multimodality in the true pushforward distribution.

BQ provides a nonlinear alternative by evaluating the conditional solver at multiple input locations. It also reports the additional outer uncertainty term Σ_k^{out} , which represents uncertainty in the numerical approximation of the outer integral under the GP model. A comparison with a classical quadrature method is given in [Section D](#). Its main limitations are the dependence on the chosen kernel and node set, and the exponential scaling of tensor-product Gauss–Hermite quadrature with the uncertain-input dimension.

Probabilistic methods cannot remove uncertainty in the IVP or prevent loss of predictability in unstable or chaotic systems. Their contribution is to separate uncertainty caused by the dynamics from uncertainty introduced by the numerical solver and the outer integration. This prevents increased numerical accuracy from being incorrectly interpreted as increased certainty about the underlying IVP.

A natural extension is an adaptive method that uses first-order propagation by default and switches to BQ when sigma-point evaluations reveal substantial nonlinear behaviour. Sparse or dimension-adaptive quadrature could improve scalability. The framework may also be useful in Bayesian inverse problems, where the forward IVP solver is repeatedly evaluated inside a likelihood or posterior computation and numerical uncertainty must be distinguished from inferential uncertainty ([Poot et al., 2026](#)).

8 Conclusion

We showed that conditional solver uncertainty, propagated IVP uncertainty, and outer quadrature uncertainty should be treated separately. The first-order method provides exact propagated moments for jointly linear–Gaussian models and an efficient approximation when the nonlinear input-to-solution map is close to affine. BQ provides a more flexible nonlinear moment transform and additionally quantifies uncertainty in the outer numerical integration.

The experiments support a practical hierarchy: first-order propagation is preferable in linear or mildly nonlinear regimes, whereas BQ is more appropriate when nonlinear input effects dominate and additional conditional solver evaluations are affordable. More generally, probabilistic methods help by identifying and separating the sources of uncertainty, rather than by eliminating uncertainty intrinsic to the IVP.

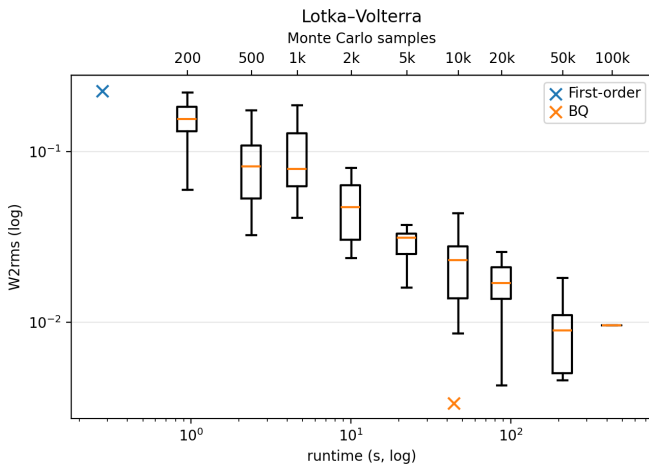


Figure 4: Lotka–Volterra IVP: distribution of MC errors (boxplots) versus runtime, measured by $W2_{rms}$ relative to a $N_{ref} = 10^6$ reference.

Acknowledgments

Tom Colemont, Frederik De Ceuster and Tjonnje G.F. Li acknowledge support by the Research Foundation – Flanders (FWO; Grant No. I000725N, I002123N, and 1178926N). Johan Suykens acknowledges support of Research Council KU Leuven C14/24/103, iBOF/23/064; Flemish Government AI Research Program.

References

- N. Bosch, P. Hennig, and F. Tronarp. Calibrated adaptive probabilistic ode solvers. In *International Conference on Artificial Intelligence and Statistics*, pages 3466–3474. PMLR, 2021.
- F.-X. Briol, C. J. Oates, M. Girolami, M. A. Osborne, and D. Sejdinovic. Probabilistic integration. *Statistical Science*, 34(1):1–22, 2019.
- T. Damour and N. Deruelle. General relativistic celestial mechanics of binary systems. i. the post-newtonian motion. In *Annales de l’IHP Physique théorique*, volume 43, pages 107–132, 1985.
- S. Daun, J. Rubin, Y. Vodovotz, and G. Clermont. Equation-based models of dynamic biological systems. *Journal of critical care*, 23(4):585–594, 2008.
- M. Gelbrich. On a formula for the l2 wasserstein metric between measures on euclidean and hilbert spaces. *Mathematische Nachrichten*, 147(1):185–203, 1990.
- D. T. Gillespie. Exact stochastic simulation of coupled chemical reactions. *The journal of physical chemistry*, 81(25):2340–2361, 1977.
- P. Hennig, M. A. Osborne, and H. P. Kersting. *Probabilistic Numerics: Computation as Machine Learning*. Cambridge University Press, 2022.
- W. O. Kermack and A. G. McKendrick. A contribution to the mathematical theory of epidemics. *Proceedings of the royal society of london. Series A, Containing papers of a mathematical and physical character*, 115(772):700–721, 1927.
- C. J. Oates and T. J. Sullivan. A modern retrospective on probabilistic numerics. *Statistics and computing*, 29(6):1335–1351, 2019.
- A. O’Hagan. Bayes–hermite quadrature. *Journal of statistical planning and inference*, 29(3):245–260, 1991.
- A. Poot, I. Rocha, P. Kerfriden, and F. van der Meer. The bayesian finite element method in inverse problems: a critical comparison between probabilistic models for discretization error. *SIAM/ASA Journal on Uncertainty Quantification*, 14(2):287–312, 2026.
- J. Prüher and M. Šimandl. Bayesian quadrature variance in sigma-point filtering. In *Informatics in Control, Automation and Robotics 12th International Conference, ICINCO 2015 Colmar, France, July 21-23, 2015 Revised Selected Papers*, pages 355–370. Springer, 2016.
- S. Särkkä and L. Svensson. *Bayesian filtering and smoothing*, volume 17. Cambridge university press, 2023.
- S. Särkkä, J. Hartikainen, L. Svensson, and F. Sandblom. On the relation between gaussian process quadratures and sigma-point methods. *arXiv preprint arXiv:1504.05994*, 2015.
- M. Schober, D. Duvenaud, and P. Hennig. Probabilistic ode solvers with runge-kutta means. volume 27, 2014.
- M. Schober, S. Särkkä, and P. Hennig. A probabilistic model for the numerical solution of initial value problems. *Statistics and Computing*, 29(1):99–122, 2019.
- A. H. Stroud, D. Secrest, et al. *Gaussian quadrature formulas*, volume 374. Prentice-Hall Englewood Cliffs, NJ, 1966.
- F. Tronarp, H. Kersting, S. Särkkä, and P. Hennig. Probabilistic solutions to ordinary differential equations as nonlinear bayesian filtering: a new perspective. *Statistics and Computing*, 29(6):1297–1315, 2019.
- D. Yao, F. Tronarp, and N. Bosch. Propagating model uncertainty through filtering-based probabilistic numerical ode solvers. 2025.

A Continuous-time jointly linear SDE

Consider the continuous-time jointly linear model

$$\frac{d}{dt}\mathbf{x}(t) = \mathbf{F}\mathbf{x}(t) + \mathbf{G}\boldsymbol{\theta} + \mathbf{L}\mathbf{w}(t), \quad t \geq t_0, \quad (66)$$

with Gaussian prior

$$\boldsymbol{\theta} \sim \mathcal{N}(\mathbf{m}_{\boldsymbol{\theta}}, \mathbf{P}_{\boldsymbol{\theta}}), \quad \mathbf{x}(t_0) \mid \boldsymbol{\theta} \sim \mathcal{N}(\mathbf{m}_{0|\boldsymbol{\theta}}, \mathbf{P}_{0|\boldsymbol{\theta}}). \quad (67)$$

Here $\mathbf{w}(t)$ denotes vector-valued white noise with spectral density $\mathbf{Q}$, that is, $\mathbb{E}[\mathbf{w}(t)\mathbf{w}(t')^\top] = \mathbf{Q}\delta(t-t')$.

A.1 Augmented-state representation

Define the augmented state

$$\mathbf{z}(t) := \begin{bmatrix} \mathbf{x}(t) \\ \boldsymbol{\theta} \end{bmatrix}. \quad (68)$$

Under the assumptions that $\boldsymbol{\theta}$ is static and that the drift is linear in $\boldsymbol{\theta}$, Eq. (66) can be written equivalently as the linear SDE

$$\begin{aligned} \frac{d}{dt}\mathbf{z}(t) &= \bar{\mathbf{F}}\mathbf{z}(t) + \bar{\mathbf{L}}\mathbf{w}(t), & \bar{\mathbf{F}} &:= \begin{bmatrix} \mathbf{F} & \mathbf{G} \\ \mathbf{0} & \mathbf{0} \end{bmatrix}, \\ \bar{\mathbf{L}} &:= \begin{bmatrix} \mathbf{L} \\ \mathbf{0} \end{bmatrix}. \end{aligned} \quad (69)$$

This augmentation is the key step that turns uncertain static parameters into part of a larger linear Gaussian dynamical system (Särkkä and Svensson, 2023).

A.2 Mean and covariance dynamics

For a time-invariant linear SDE of the form

$$\frac{d}{dt}\mathbf{x}(t) = \mathbf{F}\mathbf{x}(t) + \mathbf{b}(t) + \mathbf{L}\mathbf{w}(t), \quad \mathbf{x}(t_0) \sim \mathcal{N}(\mathbf{m}_0, \mathbf{P}_0), \quad (70)$$

the mean $\mathbf{m}(t) := \mathbb{E}[\mathbf{x}(t)]$ and covariance $\mathbf{P}(t) := \text{Cov}(\mathbf{x}(t))$ satisfy (Särkkä and Svensson, 2023)

$$\frac{d}{dt}\mathbf{m}(t) = \mathbf{F}\mathbf{m}(t) + \mathbf{b}(t), \quad (71)$$

$$\frac{d}{dt}\mathbf{P}(t) = \mathbf{F}\mathbf{P}(t) + \mathbf{P}(t)\mathbf{F}^\top + \mathbf{L}\mathbf{Q}\mathbf{L}^\top. \quad (72)$$

Specialising these equations to the augmented model Eq. (69), where $\mathbf{b}(t) \equiv 0$, the augmented mean $\mathbf{m}_z(t) := \mathbb{E}[\mathbf{z}(t)]$ satisfies

$$\frac{d}{dt}\mathbf{m}_z(t) = \bar{\mathbf{F}}\mathbf{m}_z(t), \quad \mathbf{m}_z(t) = \begin{bmatrix} \mathbf{m}_x(t) \\ \mathbf{m}_\theta(t) \end{bmatrix}, \quad (73)$$

and the augmented covariance $\mathbf{P}_z(t) := \text{Cov}(\mathbf{z}(t))$ satisfies

$$\begin{aligned} \frac{d}{dt}\mathbf{P}_z(t) &= \bar{\mathbf{F}}\mathbf{P}_z(t) + \mathbf{P}_z(t)\bar{\mathbf{F}}^\top + \bar{\mathbf{L}}\mathbf{Q}\bar{\mathbf{L}}^\top, \\ \mathbf{P}_z(t) &= \begin{bmatrix} \mathbf{P}_x(t) & \mathbf{P}_{x\theta}(t) \\ \mathbf{P}_{x\theta}(t)^\top & \mathbf{P}_\theta(t) \end{bmatrix}. \end{aligned} \quad (74)$$

Writing this equation blockwise yields the expected structure. In particular, the parameter covariance is static:

$$\frac{d}{dt}\mathbf{P}_x(t) = \mathbf{F}\mathbf{P}_x(t) + \mathbf{P}_x(t)\mathbf{F}^\top + \mathbf{G}\mathbf{P}_{x\theta}(t)^\top \quad (75a)$$

$$+ \mathbf{P}_{x\theta}(t)\mathbf{G}^\top + \mathbf{L}\mathbf{Q}\mathbf{L}^\top, \quad (75b)$$

$$\frac{d}{dt}\mathbf{P}_{x\theta}(t) = \mathbf{F}\mathbf{P}_{x\theta}(t) + \mathbf{G}\mathbf{P}_\theta(t), \quad (75c)$$

$$\frac{d}{dt}\mathbf{P}_\theta(t) = 0. \quad (75d)$$

A.3 Closed-form solutions

Let $\Sigma_0 := \mathbf{P}_z(t_0)$ denote the joint prior covariance of $(\mathbf{x}(t_0), \theta)$, and let $\mathbf{m}_z(t_0)$ be the corresponding joint prior mean. Then

$$\mathbf{m}_z(t) = \exp(\bar{\mathbf{F}}(t - t_0)) \mathbf{m}_z(t_0), \quad (76a)$$

$$\begin{aligned} \mathbf{P}_z(t) &= \exp(\bar{\mathbf{F}}(t - t_0)) \Sigma_0 \exp(\bar{\mathbf{F}}(t - t_0))^\top \\ &\quad + \int_0^{t-t_0} \exp(\bar{\mathbf{F}}s) \bar{\mathbf{L}}\mathbf{Q}\bar{\mathbf{L}}^\top \exp(\bar{\mathbf{F}}s)^\top ds. \end{aligned} \quad (76c)$$

These formulas are standard linear-SDE results (Särkkä and Svensson, 2023), but in the present context they show explicitly how parameter uncertainty is transported into state uncertainty and cross-covariance.

A.4 Discrete jointly linear SDE

The exact discretisation of Eq. (70) on a grid $\{t_k\}_{k \geq 0}$ with step size $\Delta t_k := t_{k+1} - t_k$ can be written as (Särkkä and Svensson, 2023)

$$\mathbf{x}_{k+1} = \mathbf{A}_k \mathbf{x}_k + \mathbf{u}_k + \mathbf{q}_k, \quad \mathbf{q}_k \sim \mathcal{N}(0, \mathbf{Q}_k), \quad (77a)$$

$$\mathbf{y}_k = \mathbf{C}_k \mathbf{x}_k + \mathbf{r}_k, \quad \mathbf{r}_k \sim \mathcal{N}(0, \mathbf{R}_k), \quad (77b)$$

where

$$\mathbf{A}_k = \exp(\mathbf{F} \Delta t_k), \quad (78)$$

$$\mathbf{u}_k = \int_{t_k}^{t_{k+1}} \exp(\mathbf{F}(t_{k+1} - \tau)) \mathbf{b}(\tau) d\tau, \quad (79)$$

$$\mathbf{Q}_k = \int_0^{\Delta t_k} \exp(\mathbf{F}s) \mathbf{L}\mathbf{Q}\mathbf{L}^\top \exp(\mathbf{F}s)^\top ds. \quad (80)$$

Specialising to the augmented SDE Eq. (69) with constant step size h yields

$$\mathbf{z}_{k+1} | \mathbf{z}_k \sim \mathcal{N}(\bar{\mathbf{A}}(h) \mathbf{z}_k, \bar{\mathbf{Q}}(h)), \quad (81)$$

with

$$\bar{\mathbf{A}}(h) = \exp(\bar{\mathbf{F}}h), \quad (82a)$$

$$\bar{\mathbf{Q}}(h) = \int_0^h \exp(\bar{\mathbf{F}}s) \bar{\mathbf{L}}\mathbf{Q}\bar{\mathbf{L}}^\top \exp(\bar{\mathbf{F}}s)^\top ds. \quad (82b)$$

The corresponding observation model from Section 4.1 can be written in augmented form as

$$\mathbf{y}_k = [\mathbf{C}_k \quad \mathbf{D}_k] \mathbf{z}_k + \mathbf{r}_k, \quad \mathbf{r}_k \sim \mathcal{N}(0, \mathbf{R}_k). \quad (83)$$

Lemma 1 (Block structure of $\exp(\bar{\mathbf{F}}h)$). *For any $h \geq 0$,*

$$\begin{aligned} \exp(\bar{\mathbf{F}}h) &= \begin{bmatrix} \mathbf{A}(h) & \mathbf{B}(h) \\ 0 & \mathbf{I} \end{bmatrix}, \quad \mathbf{A}(h) := \exp(\mathbf{F}h), \\ \mathbf{B}(h) &:= \left(\int_0^h \exp(\mathbf{F}s) ds \right) \mathbf{G}. \end{aligned} \quad (84)$$

Proof. Using the power-series definition of the matrix exponential,

$$\exp(\bar{\mathbf{F}}h) = \mathbf{I} + \sum_{n=1}^{\infty} \frac{h^n}{n!} \bar{\mathbf{F}}^n. \quad (85)$$

A direct induction shows that for $n \geq 1$,

$$\bar{\mathbf{F}}^n = \begin{bmatrix} \mathbf{F}^n & \mathbf{F}^{n-1}\mathbf{G} \\ 0 & 0 \end{bmatrix}. \quad (86)$$

Substituting this into the series yields

$$\exp(\bar{\mathbf{F}}h) = \begin{bmatrix} \mathbf{I} + \sum_{n=1}^{\infty} \frac{(\mathbf{F}h)^n}{n!} & \sum_{n=1}^{\infty} \frac{h^n}{n!} \mathbf{F}^{n-1}\mathbf{G} \\ 0 & \mathbf{I} \end{bmatrix}. \quad (87)$$

The top-left block equals $\exp(\mathbf{F}h) = \mathbf{A}(h)$. For the top-right block,

$$\sum_{n=1}^{\infty} \frac{h^n}{n!} \mathbf{F}^{n-1} = \sum_{n=0}^{\infty} \frac{h^{n+1}}{(n+1)!} \mathbf{F}^n \quad (88)$$

$$= \int_0^h \sum_{n=0}^{\infty} \frac{s^n}{n!} \mathbf{F}^n ds \quad (89)$$

$$= \int_0^h \exp(\mathbf{F}s) ds, \quad (90)$$

so the block equals $\mathbf{B}(h)$, proving Eq. (84). $\square$

A.5 Noise covariance for the augmented discretisation

By Lemma 1,

$$\exp(\bar{\mathbf{F}}s) \bar{\mathbf{L}} = \begin{bmatrix} \mathbf{A}(s)\mathbf{L} \\ 0 \end{bmatrix}, \quad (91)$$

and hence Eq. (82) implies

$$\begin{aligned} \bar{\mathbf{Q}}(h) &= \begin{bmatrix} \mathbf{Q}(h) & 0 \\ 0 & 0 \end{bmatrix}, \\ \mathbf{Q}(h) &= \int_0^h \mathbf{A}(s) \mathbf{L} \mathbf{Q} \mathbf{L}^\top \mathbf{A}(s)^\top ds. \end{aligned} \quad (92)$$

Consequently, the discrete augmented dynamics Eq. (81) can be written as

$$\begin{aligned} \mathbf{x}_{k+1} &= \mathbf{A}(h) \mathbf{x}_k + \mathbf{B}(h) \boldsymbol{\theta} + \mathbf{q}_k, & \mathbf{q}_k &\sim \mathcal{N}(0, \mathbf{Q}(h)), \\ \boldsymbol{\theta}_{k+1} &= \boldsymbol{\theta}_k. \end{aligned} \quad (93)$$

A.6 Example: IWP(1)

For the integrated Wiener process prior of order 1 (IWP(1)) with diffusion intensity q_c ,

$$\mathbf{F} = \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix}, \quad \mathbf{L} = \begin{bmatrix} 0 \\ 1 \end{bmatrix}, \quad \mathbf{Q} = q_c, \quad (94)$$

and Eqs. (84) and (92) yield the familiar closed forms

$$\mathbf{A}(h) = \begin{bmatrix} 1 & h \\ 0 & 1 \end{bmatrix}, \quad \mathbf{Q}(h) = q_c \begin{bmatrix} h^3/3 & h^2/2 \\ h^2/2 & h \end{bmatrix}. \quad (95)$$

B Kernel used for Bayesian quadrature

Here, we make explicit which kernel is used in the BQ construction of Section 5.1, and how the quantities $\mathbf{K}, \boldsymbol{\ell}, \mathbf{L}, \mathbf{W}$ and σ_{int}^2 are computed in closed form.

The construction in Section 5.1 is written for a generic scalar kernel $k(\cdot, \cdot)$. In the present work, this kernel is instantiated as the exponentiated quadratic (EQ) kernel; see, e.g., (O’Hagan, 1991; Pr  her and Šimandl, 2016).

Whitened coordinates. Recall from Eq. (50) that the uncertain input is written as

$$\begin{aligned} \mathbf{u} &= \mathbf{m}_u + \mathbf{L}_u \boldsymbol{\xi}, & \boldsymbol{\xi} &\sim \mathcal{N}(0, \mathbf{I}_{d_u}), \\ \mathbf{L}_u \mathbf{L}_u^\top &= \mathbf{P}_u. \end{aligned} \quad (96)$$

Kernel. The scalar GP kernel used is

$$k(\boldsymbol{\xi}, \boldsymbol{\xi}') = \alpha^2 \exp\left(-\frac{\|\boldsymbol{\xi} - \boldsymbol{\xi}'\|^2}{2\lambda^2}\right), \quad (97)$$

with amplitude $\alpha^2 > 0$ and length-scale $\lambda > 0$.

Gauss–Hermite sigma points. As described in Section 5.1, the BQ construction itself does not prescribe a unique sigma-point set. In this paper, the points $\{\boldsymbol{\xi}^{(i)}\}_{i=1}^N$ are chosen as tensor-product Gauss–Hermite points for the Gaussian measure; see also (Stroud et al., 1966).

In one dimension, the Gauss–Hermite rule of order n approximates integrals of the form

$$\int_{-\infty}^{\infty} h(x) e^{-x^2} dx \approx \sum_{j=1}^n \omega_j h(x_j), \quad (98)$$

where x_j are the Hermite nodes and ω_j the corresponding weights. Since in Eq. (50) the reference measure is standard Gaussian, we rewrite

$$\begin{aligned} \mathbb{E}_{\boldsymbol{\xi}}[\varphi(\boldsymbol{\xi})] &= \int_{\mathbb{R}} \varphi(\xi) \frac{1}{\sqrt{2\pi}} e^{-\xi^2/2} d\xi \\ &= \frac{1}{\sqrt{\pi}} \int_{\mathbb{R}} \varphi(\sqrt{2}x) e^{-x^2} dx \\ &\approx \sum_{j=1}^n \bar{\omega}_j \varphi(\bar{\xi}_j), \end{aligned} \quad (99)$$

with rescaled nodes and weights

$$\bar{\xi}_j := \sqrt{2} x_j, \quad \bar{\omega}_j := \frac{\omega_j}{\sqrt{\pi}}. \quad (100)$$

In d_u dimensions, the tensor-product rule is obtained by taking Cartesian products of the one-dimensional nodes and multiplying the corresponding weights. Thus,

$$\boldsymbol{\xi}^{(i_1, \dots, i_{d_u})} = \sqrt{2} \begin{bmatrix} x_{i_1} \\ \vdots \\ x_{i_{d_u}} \end{bmatrix}, \quad (101)$$

with tensor-product weights

$$\bar{w}_{i_1, \dots, i_{d_u}} = \prod_{r=1}^{d_u} \frac{\omega_{i_r}}{\sqrt{\pi}}. \quad (102)$$

These are the sigma points used to evaluate the conditional solver mean map g_k from Eq. (51).

B.1 Closed-form kernel quantities

Given the sigma points $\{\boldsymbol{\xi}^{(i)}\}_{i=1}^N$, we now define the matrices and vectors appearing in Eqs. (53), (54), (57) and (60).

Gram matrix. The Gram matrix from Eq. (53) is

$$\mathbf{K}_{ij} = k(\boldsymbol{\xi}^{(i)}, \boldsymbol{\xi}^{(j)}) = \alpha^2 \exp\left(-\frac{\|\boldsymbol{\xi}^{(i)} - \boldsymbol{\xi}^{(j)}\|^2}{2\lambda^2}\right). \quad (103)$$

Kernel mean vector. The vector $\boldsymbol{\ell}$ from Eq. (54) has entries

$$\ell_i := \mathbb{E}_{\boldsymbol{\xi}}[k(\boldsymbol{\xi}, \boldsymbol{\xi}^{(i)})], \quad \boldsymbol{\xi} \sim \mathcal{N}(0, \mathbf{I}_{d_u}). \quad (104)$$

For the EQ kernel Eq. (97), this Gaussian integral is available in closed form:

$$\ell_i = \alpha^2 \left(\frac{\lambda^2}{\lambda^2 + 1}\right)^{d_u/2} \exp\left(-\frac{\|\boldsymbol{\xi}^{(i)}\|^2}{2(\lambda^2 + 1)}\right). \quad (105)$$

Kernel second-moment matrix. The matrix $\mathbf{L}$ from Eq. (55) has entries

$$\mathbf{L}_{ij} := \mathbb{E}_{\boldsymbol{\xi}}[k(\boldsymbol{\xi}, \boldsymbol{\xi}^{(i)}) k(\boldsymbol{\xi}, \boldsymbol{\xi}^{(j)})]. \quad (106)$$

Again, for the EQ kernel this integral can be evaluated in closed form:

$$\begin{aligned} \mathbf{L}_{ij} = & \alpha^4 \left(\frac{\lambda^2}{\lambda^2 + 2} \right)^{d_u/2} \\ & \times \exp \left(\frac{2(\boldsymbol{\xi}^{(i)})^\top \boldsymbol{\xi}^{(j)} - (\lambda^2 + 1)(\|\boldsymbol{\xi}^{(i)}\|^2 + \|\boldsymbol{\xi}^{(j)}\|^2)}{2\lambda^2(\lambda^2 + 2)} \right). \end{aligned} \quad (107)$$

Bayesian quadrature weights. With $\mathbf{K}$, $\boldsymbol{\ell}$, and $\mathbf{L}$ defined above, the BQ moment transform uses

$$\mathbf{w} = \mathbf{K}^{-1}\boldsymbol{\ell}, \quad \mathbf{W} = \mathbf{K}^{-1}\mathbf{L}\mathbf{K}^{-1}, \quad (108)$$

exactly as in Eqs. (56) and (57). Therefore, the Bayesian quadrature approximation of the propagated mean is

$$\boldsymbol{\mu}_k^{\text{BQ}} = \mathbf{G}_k^\top \mathbf{w}, \quad (109)$$

with $\mathbf{G}_k$ defined in Section 5.1 by

$$(\mathbf{G}_k)_{i,:} = g_k(\boldsymbol{\xi}^{(i)})^\top. \quad (110)$$

Integration-variance term. Finally, the additional scalar variance term from Eq. (60) is

$$\sigma_{\text{int}}^2 = \bar{k} - \text{tr}(\mathbf{K}^{-1}\mathbf{L}). \quad (111)$$

For the EQ kernel Eq. (97),

$$k(\boldsymbol{\xi}, \boldsymbol{\xi}) = \alpha^2 \quad \text{for all } \boldsymbol{\xi}, \quad (112)$$

so that

$$\bar{k} = \mathbb{E}_{\boldsymbol{\xi}}[k(\boldsymbol{\xi}, \boldsymbol{\xi})] = \alpha^2. \quad (113)$$

Hence

$$\sigma_{\text{int}}^2 = \alpha^2 - \text{tr}(\mathbf{K}^{-1}\mathbf{L}). \quad (114)$$

B.2 Resulting BQ moment approximation

Using the notation of Section 5.1, the covariance contribution induced by uncertain IVP inputs is approximated as

$$\hat{\boldsymbol{\Pi}}_k^{\text{BQ}} = \mathbf{G}_k^\top \mathbf{W} \mathbf{G}_k - \boldsymbol{\mu}_k^{\text{BQ}} (\boldsymbol{\mu}_k^{\text{BQ}})^\top, \quad (115)$$

which is exactly the structure of Eq. (59). Together with the Gauss–Hermite approximation of the average conditional covariance,

$$\mathbb{E}_{\mathbf{u}}[\widehat{\mathbf{P}}_k(\mathbf{u})] = \sum_{i=1}^N w_i \bar{\mathbf{P}}_k(\mathbf{u}^{(i)}), \quad (116)$$

this yields the propagated covariance approximation

$$\mathbf{P}_k^{\text{prop,BQ}} = \mathbb{E}_{\mathbf{u}}[\widehat{\mathbf{P}}_k(\mathbf{u})] + \hat{\boldsymbol{\Pi}}_k^{\text{BQ}} + \sigma_{\text{int}}^2 \mathbf{I}_{d_x}. \quad (117)$$

C Used parameters for benchmark problems

In Table 2, the settings for the benchmark problems are given.

D Matched-node comparison with classical quadrature

To compare the outer integration rules directly, classical Gauss–Hermite quadrature and BQ are evaluated on the same tensor-product Gauss–Hermite node set. At each node, the same conditional probabilistic ODE solver is run. Hence the conditional solver output is fixed across the two methods, and the comparison isolates the effect of the outer moment weights.

Classical Gauss–Hermite quadrature uses deterministic weights $\bar{\mathbf{w}}$. With $\mathbf{G}_k \in \mathbb{R}^{N \times d_x}$ denoting the matrix of conditional solver means at time t_k , its propagated mean and IVP covariance are

$$\boldsymbol{\mu}_k^{\text{GH}} = \mathbf{G}_k^\top \bar{\mathbf{w}}, \quad (118)$$

$$\hat{\boldsymbol{\Pi}}_k^{\text{GH}} = \mathbf{G}_k^\top \text{diag}(\bar{\mathbf{w}}) \mathbf{G}_k - \boldsymbol{\mu}_k^{\text{GH}} (\boldsymbol{\mu}_k^{\text{GH}})^\top. \quad (119)$$

BQ uses the same evaluations $\mathbf{G}_k$, but replaces the deterministic Gauss–Hermite weights by kernel-dependent Bayesian quadrature weights,

$$\mathbf{w} = \mathbf{K}^{-1}\boldsymbol{\ell}, \quad \mathbf{W} = \mathbf{K}^{-1}\mathbf{L}\mathbf{K}^{-1}. \quad (120)$$

The corresponding moments are

$$\boldsymbol{\mu}_k^{\text{BQ}} = \mathbf{G}_k^\top \mathbf{w}, \quad (121)$$

$$\hat{\boldsymbol{\Pi}}_k^{\text{BQ}} = \mathbf{G}_k^\top \mathbf{W} \mathbf{G}_k - \boldsymbol{\mu}_k^{\text{BQ}} (\boldsymbol{\mu}_k^{\text{BQ}})^\top. \quad (122)$$

Thus, in Fig. 5, the difference between GH and BQ is not caused by different solver calls or different sigma-point locations, but by the different weights used to approximate the outer Gaussian-weighted moments. Both methods converge to the same result as the number of nodes increases. For GH, the convergence rate depends primarily on how accurately the propagated quantity can be represented by the chosen quadrature order, whereas for BQ it additionally depends on the kernel and its hyperparameters.

The second distinction is that BQ also reports an outer quadrature uncertainty,

$$\boldsymbol{\Sigma}_k^{\text{out}} = \sigma_{\text{int}}^2 \mathbf{I}_{d_x}, \quad \sigma_{\text{int}}^2 = \bar{k} - \text{tr}(\mathbf{K}^{-1}\mathbf{L}),$$

whereas deterministic Gauss–Hermite quadrature does not quantify an uncertainty estimate. The main benefit of BQ in this comparison is therefore not necessarily a smaller matched-node point estimate, but the availability of an explicit uncertainty estimate for the outer integration step.

E Trajectory-level benchmark diagnostics

Fig. 6 gives trajectories for the three benchmark problems considered in Section 6.

The linear benchmark in Fig. 6a illustrates the expected expansion of initial-condition uncertainty. Since the exact flow is $x(t) = e^t x(0)$, the variance induced by the uncertain initial condition grows like $e^{2t} \text{Var}(x(0))$. The

Table 2: Benchmark IVPs and uncertainty specifications. The Gauss–Hermite (GH) settings are those used to evaluate the BQ moments.

Problem	Uncertain inputs $\mathbf{u}$	Prior on $\mathbf{u}$	d_u	GH order	# nodes N
Linear	$x(0)$	$\mathcal{N}(1, 10^{-2})$	1	20	20
Logistic	b	$\mathcal{N}(3, 10^{-2})$	1	50	50
Lotka–Volterra	$\mathbf{x}(0)$	$\mathcal{N}([5, 5]^\top, 0.3\mathbf{I}_2)$	2	12	$12^2 = 144$

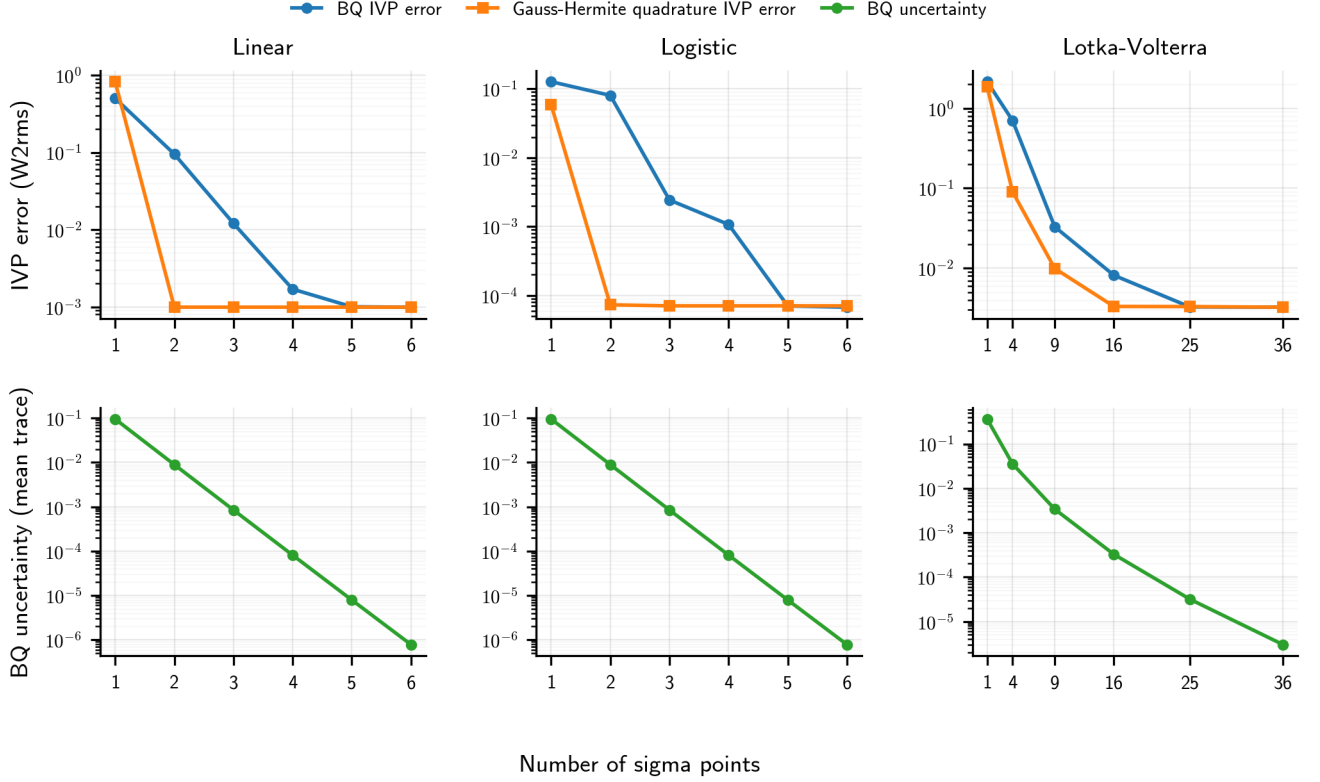


Figure 5: Matched-node comparison of classical Gauss–Hermite quadrature and BQ. Both methods use the same tensor-product Gauss–Hermite nodes and the same conditional probabilistic ODE solver evaluations. The difference is only in the outer integration weights. The top row reports the $W_{2,rms}$ error of the propagated IVP moments relative to a refined reference. The bottom row reports the mean trace of the BQ outer uncertainty Σ_k^{out} .

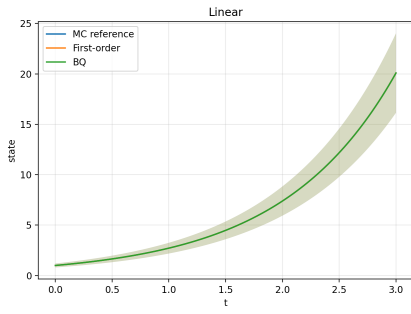
propagated uncertainty therefore increases over time and is not reduced by making the numerical solve more accurate. This is the behaviour that the separation between conditional solver uncertainty and IVP uncertainty is intended to preserve.

The logistic benchmark in Fig. 6b shows a different qualitative regime. The dynamics approach an attracting level, and the propagated uncertainty does not grow in the same way as in the unstable linear problem. In this experiment the uncertain input is the carrying-capacity parameter b , rather than the initial condition. The uncertainty therefore affects the asymptotic level of the trajectory. This illustrates that initial-condition and parameter uncertainty are treated within the same outer marginalisation framework, but their propagated effects differ through the corresponding input-to-solution sensitivities.

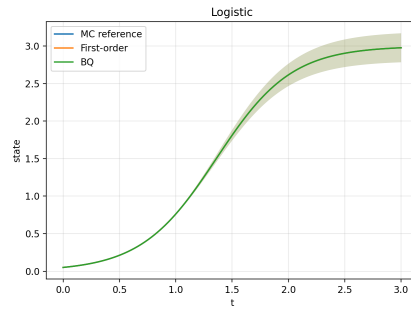
The Lotka–Volterra benchmark in Fig. 6c illustrates the nonlinear coupled case. Here the uncertain initial condition is propagated through oscillatory dynamics, so the input-to-solution map is no longer close to affine. This explains why the first-order approximation is less

accurate in this benchmark, while BQ benefits from evaluating the nonlinear conditional solver at multiple input locations.

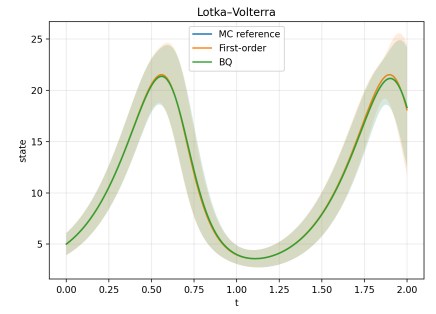
Overall, these trajectory plots show that the proposed propagation framework does not impose either expansion or contraction of uncertainty. The qualitative behaviour is determined by the IVP dynamics and by the uncertain input under consideration.



(a) Linear benchmark.



(b) Logistic benchmark.



(c) Lotka-Volterra benchmark.

Figure 6: Trajectory-level propagated uncertainty bands for the three benchmark IVPs. The blue curve denotes the Monte Carlo reference, the orange curve the first-order approximation, and the green curve the BQ approximation. Shaded regions show the corresponding propagated uncertainty 95% bands.