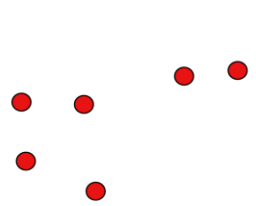


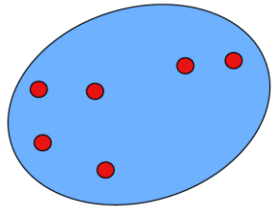
# Learning Where to Learn

Nicolas Guerra

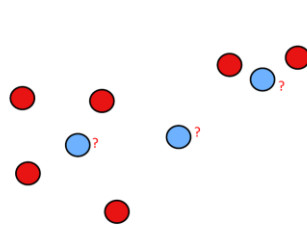
Goal: Find optimal  $\nu$  such that  $\hat{F}$  performs well on unseen distributions.  
( $\hat{F}$  learned  $F : X \rightarrow Y$  using data  $x \sim \nu$ )



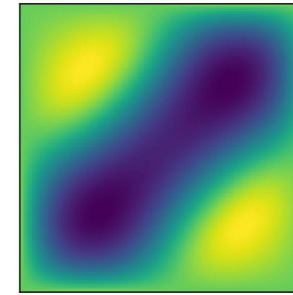
Family of test distributions.



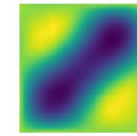
Ideally, train on a large enough distribution (blue).



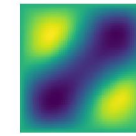
In practice, constrained.



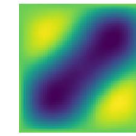
true conductivity



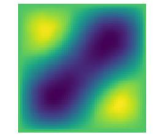
$\nu_1$



$\nu_2$



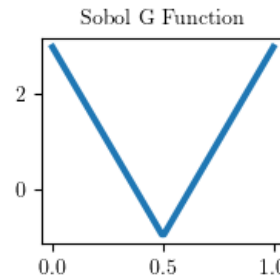
$\nu_3$



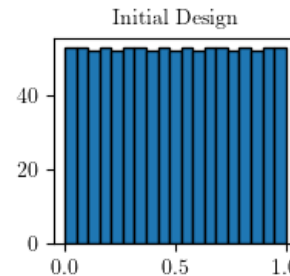
$\nu_4$

Ways to represent  $\nu$ :

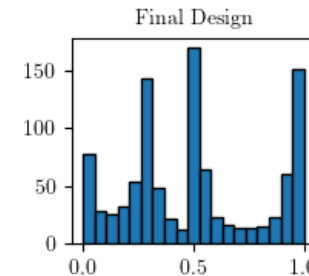
- Parametric, e.g. Gaussian process
- Nonparametric, i.e.  $\nu = \sum_{i=1}^N w_i \delta_{x_i}$



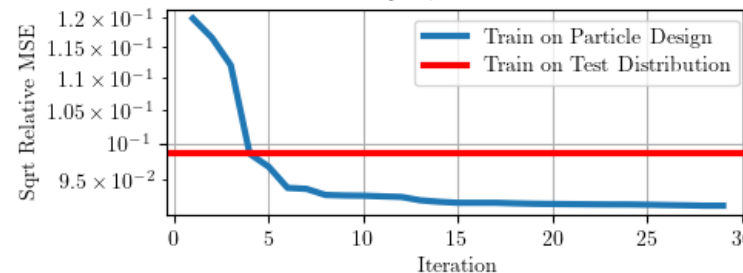
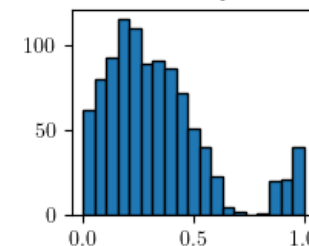
Step 30, Loss: 0.0914



Iteration



Test Samples



Applications/Inverse Problems:

- Electrical Impedance Tomography
- Darcy Flow
- Radiative Transfer



Cornell University

# CGKN: Conditional Gaussian Koopman Network

## A Stochastic Digital Twin for Modeling, Forecasting, Data Assimilation, and Uncertainty Quantification

Nan Chen<sup>1\*</sup>, Chuanqi Chen<sup>2</sup>, Zhongrui Wang<sup>1</sup> and Jinlong Wu<sup>2</sup>

<sup>1</sup>UW-Madison (Math) <sup>2</sup>UW-Madison (ME) \*chennan@math.wisc.edu

Consider a spatiotemporal discretization of a general complex turbulent system (e.g., PDEs). In many cases, only the measurements of part of the state variables are available (e.g., sparse obs):

$$\begin{aligned} \mathbf{u}_1(n+1) &= \mathcal{G}_1(\mathbf{u}_1(n), \mathbf{u}_2(n)), & \xrightarrow{\text{CGKN}} \quad \mathbf{u}_1(n+1) &= \mathbf{F}_1(\mathbf{u}_1(n)) + \mathbf{G}_1(\mathbf{u}_1(n))\mathbf{v}(n) + \sigma_1\epsilon_1(n+1), \\ \mathbf{u}_2(n+1) &= \mathcal{G}_2(\mathbf{u}_1(n), \mathbf{u}_2(n)), & \mathbf{v}(n+1) &= \mathbf{F}_2(\mathbf{u}_1(n)) + \mathbf{G}_2(\mathbf{u}_1(n))\mathbf{v}(n) + \sigma_2\epsilon_2(n+1), \end{aligned}$$

where  $\mathbf{v}(n) = \varphi(\mathbf{u}_2(n))$ , while  $\mathbf{F}_1$ ,  $\mathbf{F}_2$ ,  $\mathbf{G}_1$ , and  $\mathbf{G}_2$  contain neural operators.

The CGKN transforms **general nonlinear systems** into **highly nonlinear neural differential equations with conditional Gaussian structures** (Chen et. al., *JCP* 2024, *CMAME* 2025).

- **Physics + Machine Learning.** The mathematical justification for the conditional Gaussian nonlinear stochastic systems has been demonstrated in (Chen & Majda 2018). Many well-known complex dynamical systems belong to this framework (Lorenz, Boussinesq, etc).
- **Data Assimilation.** It aims to retain essential nonlinear components while applying systematic simplifications to facilitate the development of **analytic formulae for nonlinear DA**. This allows for **incorporating the DA performance into the deep learning training**.
- **Uncertainty Quantification.** Even though the distribution in the latent space is conditional Gaussian, when mapping back to physical space via the decoder, **the posterior distribution can be highly non-Gaussian**.

## Test Example:

2D Navier-Stokes Equation

Resolution:  $128 \times 128$

Observation:  $8 \times 8$

| Methods \ Examples | Navier-Stokes Equations |                |
|--------------------|-------------------------|----------------|
|                    | DA Error                | Forecast Error |
| CGKN               | 6.0940e+01              | 1.9754e+01     |
| EnKF               | 6.9010e+01              | —              |
| Interpolation      | 1.2844e+02              | —              |
| DNN                | —                       | 1.0936e+02     |
| CNN                | —                       | 3.0600e+01     |
| FNO                | —                       | 1.7129e+01     |

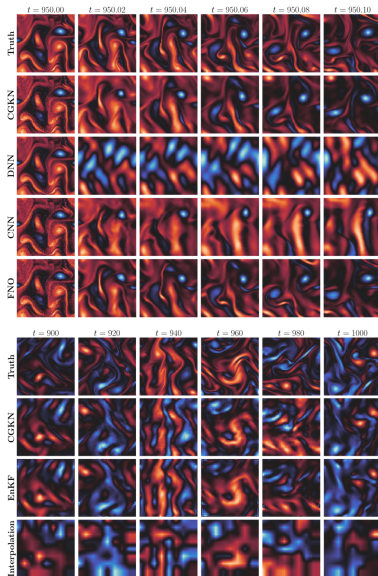
(All NNs have similar numbers of hyperparameters.)

Right top: Forecast

Right bottom: Data assimilation

**Computational efficiency** compared to EnKF:

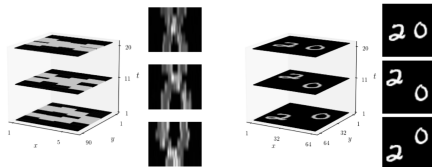
- 600x faster for Viscous Burgers Equation
- 125x faster for Kuramoto-Sivashinsky Equation
- 300x faster for Navier-Stokes Equation



CGKN has been adopted to recover the ocean field and topography (sea mountains), modeling equatorial Pacific Earth systems, conduct nonlinear Lagrangian data assimilation, incorporate memory in the system (Mori-Zwanzig formalism). An online model correction algorithm has also been developed (Chen, et. al., 2023, *Chaos*; Wang, et. al., 2025).

# Motivation & Problem Setup

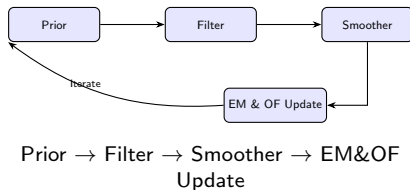
- **Dynamic X-ray CT** seeks to reconstruct a sequence of images over time from limited projection data.
- **Digital twins** rely on accurate and efficient data assimilation to monitor and predict complex dynamical systems, making scalable sequential reconstruction methods essential for real-time applications.
- Some problems are **ill-posed, high-dimensional**, and subject to **radiation exposure limits**:
  - Small measurement noise can cause large reconstruction errors.
  - Only a limited number of projections can be collected safely.
  - Traditional “all-at-once” methods scale poorly with time and memory.
- **Goal:** Develop a *scalable, adaptive framework* that processes data sequentially and doesn't require major parameter tuning, memory, or time.





# Sequential Filtering Framework

- Employs a **Kalman filtering approach** for sequential estimation in time-dependent imaging.
- Uses a **reduced-order model** for scalability.
- Core components:
  - **Kalman Filter & RTS Smoother** — dynamic estimation and temporal refinement
  - **Expectation–Maximization (EM)** — learns noise statistics automatically
  - **Optical Flow (OF)** — estimates and refines motion between frames

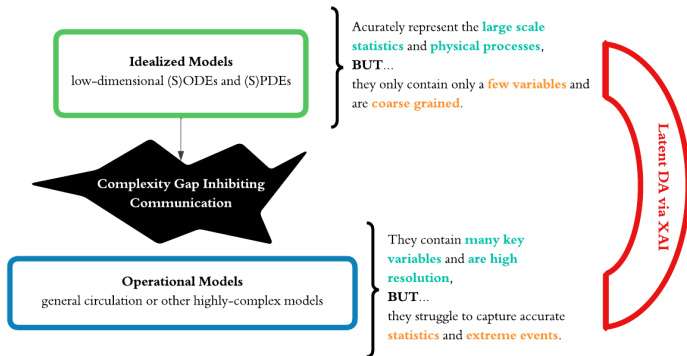




# Bridging Idealized and Operational Models: An Explainable AI Framework for Earth System Emulators

Pouria Behnoudfar<sup>1</sup>, Charlotte Moser<sup>1\*</sup>, Marc Bocquet<sup>2</sup>, Sibö Cheng<sup>2</sup>, Nan Chen<sup>1</sup>

\* crmoser2@wisc.edu



## Results of implementation

The framework produces a **bridging model** that can efficiently generate numerous high-quality synthetic datasets that **mimic nature**, facilitate the study of **extreme events**, and **provide training datasets** for many machine learning tasks, including building effective **digital twins**.

- The framework is highly efficient, featuring a **reconfigured latent data assimilation** technique specifically designed to deal with large dimension discrepancies between **idealized** and **operational** models.
- The resulting **computationally efficient bridging model** is high resolution like the **operational model** but is dynamically and statistically improved through assimilation with pseudo observations of the **idealized model**.
- The process is explainable, allowing us to trace how specific features from the **idealized model** correct specific biases in the **operational model**.

# Implementation for Representing the Dominant Interannual Phenomena: ENSO

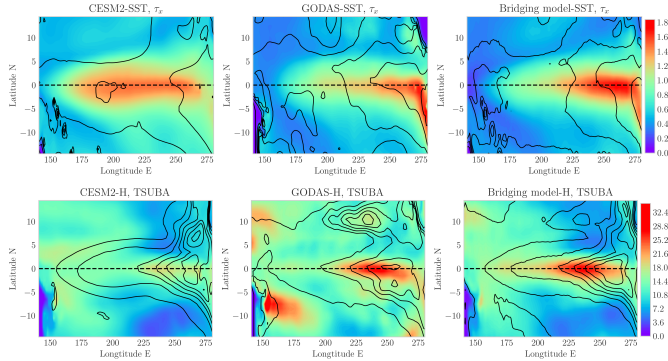


Figure: Standard deviation of different physical fields.

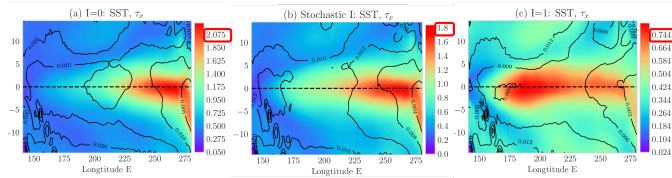


Figure: What-if scenarios for strengthened and weakened Walker circulation.

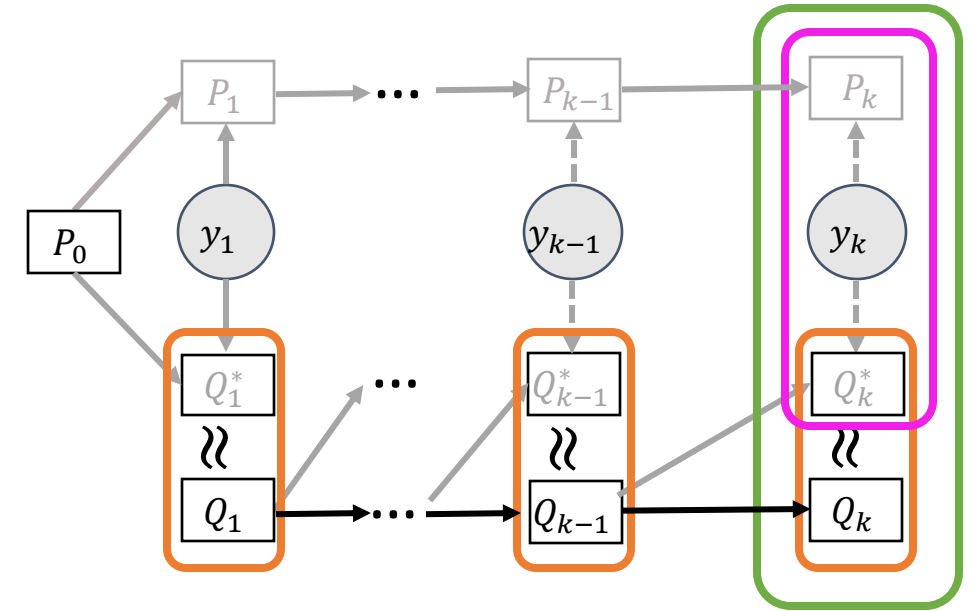
- The framework corrects both on equator and off equator biases in the **operational model** for both **pseudo-observed** and unobserved variables.
- By manipulating the long term background forcing in the **idealized model** the **bridging model** can effectively produce high resolution representations of different scenarios that mimic observed decadal variability.
- The **bridging model** provides a new way to build physics-assisted digital twins, which can efficiently test what-if scenarios.

# Towards understanding the errors in online Bayesian data assimilation

Liliang Wang and Alex Gorodetsky, University of Michigan

## Background & Motivation

- Many real-world tasks require assimilating data incrementally for real-time prediction and decision making :
  - Robotics navigation
  - Dynamic soaring of UAVs
- Practical **Online** Bayesian data assimilation methods are based on approximations: true posterior  $P_k$  is approximated with  $Q_k$
- Gap between analysis theory development and methodology development
  - The majority of the existing theoretical work are:
    - Asymptotic results
    - Developed for a specific method
  - Non-asymptotic** analysis theory for **general** online Bayesian data assimilation methods is needed



→ : prior-to-posterior map

- learning error  $d(P_k, Q_k)$  comes from two sources:
    - Incremental approximation error  $d(Q_k^*, Q_k)$
    - prior perturbation error  $d(P_k, Q_k^*)$
- $d(Q_k^*, Q_k)$ : fairly studied     $d(P_k, Q_k^*)$ : **not well studied**

# Main Results <sup>[1]</sup>

- The prior perturbation error  $d(P_k, Q_k^*)$  is caused by the prior error  $d(P_{k-1}, Q_{k-1})$ . What is the relation between the two?

**Pointwise global Lipschitz prior-to-posterior stability :**

$\forall P_{k-1}$ , we have

$$d(P_k, Q_k^*) \leq K(y_k, P_{k-1}) d(P_{k-1}, Q_{k-1}), \quad \forall Q_{k-1} \text{ (global)}$$

Holds for  $d$  being

- total variation distance
- Hellinger distance
- Wasserstein-1 distance

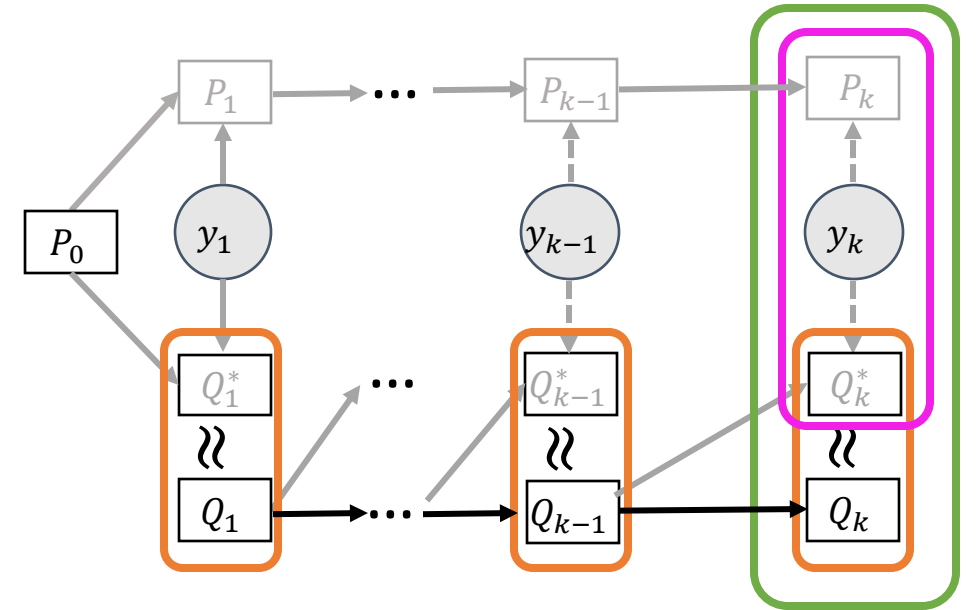
Holds in three contexts

- inverse problems
- state estimation
- joint state-parameter estimation

- Two **upper bounds** on the learning error  $d(P_k, Q_k)$ :

$$d(P_k, Q_k) \leq \sum_{j=1}^{k-1} c_1(y_k, P_{j:k-1}) d(Q_j^*, Q_j) + d(Q_k^*, Q_k)$$

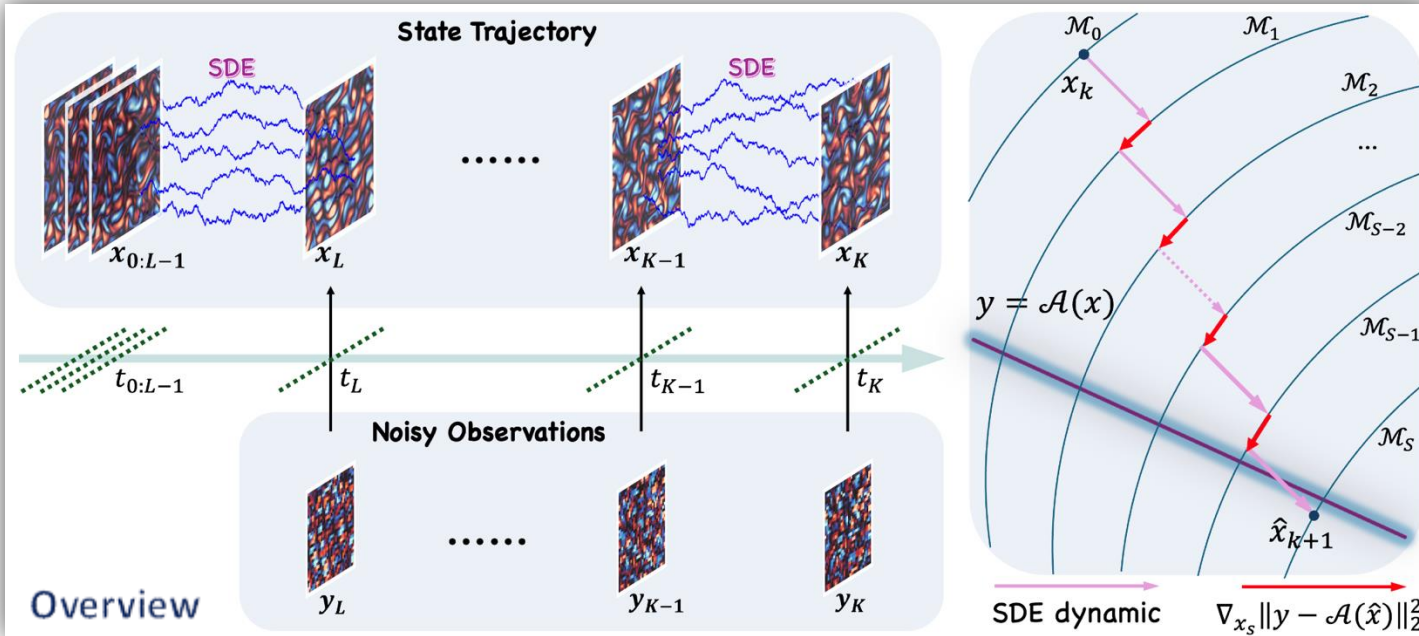
$$d(P_k, Q_k) \leq \sum_{j=1}^{k-1} c_2(y_k, Q_{j:k-1}) d(Q_j^*, Q_j) + d(Q_k^*, Q_k)$$



$\longrightarrow$  : prior-to-posterior map

- Sufficient conditions for **learning error decay**, i.e.,  
 $d(P_k, Q_k) \leq d(P_{k-1}, Q_{k-1})$

# FlowDAS: A Stochastic Interpolants-based Framework for Data Assimilation



- ❖ Consider a stochastic process  $\mathbf{X}_s$  defined over the interval  $s \in [0, 1]$ .
  - ❖ Initial state:  $\mathbf{X}_0 \sim \pi(\mathbf{X}_0)$
  - ❖ Final state:  $\mathbf{X}_1 \sim q(\mathbf{X}_1 | \mathbf{X}_0)$
- ❖ A *Stochastic Interpolant* can be described as:
  - $\mathbf{I}_s = \alpha_s \mathbf{X}_0 + \beta_s \mathbf{X}_1 + \sigma_s \mathbf{W}_s$
  - ❖  $\mathbf{W}_s$  is a Wiener process. The time-varying coefficients are defined to satisfy with boundary conditions:  $\alpha_s = 1 - s$ ,  $\beta_s = s^2$ ,  $\sigma_s = 1 - s$ .
- ❖ The velocity of the interpolant path,  $\mathbf{R}_s$ , is given by:
  - $\mathbf{R}_s = \dot{\alpha}_s \mathbf{X}_0 + \dot{\beta}_s \mathbf{X}_1 + \dot{\sigma}_s \mathbf{W}_s$
- ❖ Consider the following SDE:
  - $d\mathbf{X}_s = \mathbf{b}_s(\mathbf{X}_s, \mathbf{X}_0) ds + \sigma_s d\mathbf{W}_s$
- ❖ The drift term  $\mathbf{b}_s(\mathbf{X}_s, \mathbf{X}_0)$  is defined as the minimizer of the cost function:
  - $\mathcal{L}_b(\hat{\mathbf{b}}_s) = \int_0^1 \mathbb{E} [\|\hat{\mathbf{b}}_s(\mathbf{I}_s, \mathbf{X}_0) - \mathbf{R}_s\|^2] ds$
- ❖ Then, one can prove that:  $\text{Law}(\mathbf{I}_s | \mathbf{X}_0) = \text{Law}(\mathbf{X}_s)$ .

**Predictor**

- ❖ Derive the observation-aware drift function:

$$\begin{aligned} \mathbf{b}_s(\mathbf{X}_s, \mathbf{y}, \mathbf{X}_0) &= \mathbf{b}_s(\mathbf{X}_s, \mathbf{X}_0) + \frac{\nabla \log p(\mathbf{y} | \mathbf{X}_s, \mathbf{X}_0)}{\lambda_s \beta_s} \\ &= \mathbf{b}_s(\mathbf{X}_s, \mathbf{X}_0) + \frac{\sum_{j=1}^J \omega_j \nabla \log p(\mathbf{y} | \mathbf{X}_1^{(j)})}{\lambda_s \beta_s} \end{aligned}$$

- ❖ How to estimate the posterior likelihood  $\nabla \log p(\mathbf{y} | \mathbf{X}_1^{(j)})$ ?

- ❖ Acceleration of  $\mathbf{X}_1$  estimation:

- ❖ First-order Euler-Milstein method:

$$\hat{\mathbf{X}}_1 = \mathbf{b}_s(\mathbf{X}_s, \mathbf{X}_0)(1 - s) + \int \sigma_s d\mathbf{W}_s$$

- ❖ Second-order stochastic Runge-Kutta method:

$$\hat{\mathbf{X}}_1' = \frac{\mathbf{b}_s(\mathbf{X}_s, \mathbf{X}_0) + \mathbf{b}_1(\hat{\mathbf{X}}_1, \mathbf{X}_0)}{2}(1 - s) + \int \sigma_s d\mathbf{W}_s$$

**Corrector**

$$\begin{aligned} \tilde{\mathbf{I}}_s^k &= \alpha_s \mathbf{x}_k + \beta_s \mathbf{x}_{k+1} + \sqrt{s} \sigma_s \mathbf{z}_k \\ \mathbf{R}_s^k &= \dot{\alpha}_s \mathbf{x}_k + \dot{\beta}_s \mathbf{x}_{k+1} + \sqrt{s} \dot{\sigma}_s \mathbf{z}_k, \\ \mathcal{L}_b^{\text{emp}}(\hat{\mathbf{b}}) &= \frac{1}{K'} \sum_{k \in \mathcal{B}_{K'}} \int_0^1 \ell(\hat{\mathbf{b}}_s(\tilde{\mathbf{I}}_s^k, \mathbf{x}_k), \mathbf{R}_s^k) ds, \end{aligned}$$

## Algorithm 1 Training

- 1: **Input:** Dataset  $\mathbf{x}_{0:K}$ ; minibatch size  $K' \leq K$ ; coefficients  $\alpha_s, \beta_s, \sigma_s$
- 2: **repeat**
- 3:   Compute the empirical loss  $\mathcal{L}_b^K[\hat{\mathbf{b}}]$  in Eq. 14
- 4:   Take the gradient step on  $\mathcal{L}_b^K[\hat{\mathbf{b}}]$  to update  $\hat{\mathbf{b}}_s$
- 5: **until** converged
- 6: **return** drifts  $\hat{\mathbf{b}}_s$

## Algorithm 2 Sampling

- 1: **Input:** Observation  $\mathbf{y}_{1:K}$ , the measurement map  $\mathcal{A}$ , initial state  $\mathbf{x}_0$ , model  $\hat{\mathbf{b}}_s(\mathbf{X}, \mathbf{X}_0)$ , noise coefficient  $\sigma_s$ , grid  $s_0 = 0 < s_1 < \dots < s_N = 1$ , i.i.d.  $\mathbf{z}_n \sim \mathcal{N}(0, \mathbf{I}_D)$  for  $n = 0 : N - 1$ , step size  $\zeta_n$ , Monte Carlo sampling times  $J$
- 2: Set  $\hat{\mathbf{x}}_0 \leftarrow \mathbf{x}_0$
- 3: Set the  $(\Delta s)_n = s_{n+1} - s_n, n = 0 : N - 1$
- 4: **for**  $k = 0$  to  $K - 1$  **do**
- 5:    $\mathbf{X}_{s_0}, \mathbf{y} \leftarrow \hat{\mathbf{x}}_k, \mathbf{y}_{k+1}$
- 6:   **for**  $n = 0$  to  $N - 1$  **do**
- 7:      $\mathbf{X}'_{s_{n+1}} = \mathbf{X}_{s_n} + \hat{\mathbf{b}}_s(\mathbf{X}_{s_n}, \mathbf{X}_{s_0})(\Delta s)_n + \sigma_{s_n} \sqrt{(\Delta s)_n} \mathbf{z}_n$
- 8:      $\{\hat{\mathbf{X}}_1^{(j)}\}_{j=1}^J \leftarrow \text{Posterior estimation}(\hat{\mathbf{b}}_s, s_n, \mathbf{X}_0, \mathbf{X}_{s_n})$
- 9:      $\{w_j\}_{j=1}^J \leftarrow \text{Softmax function}(\{\|\mathbf{y} - \mathcal{A}(\hat{\mathbf{X}}_1^{(j)})\|_2^2\}_{j=1}^J)$
- 10:      $\mathbf{X}_{s_{n+1}} = \mathbf{X}'_{s_{n+1}} - \zeta_n \nabla_{\mathbf{X}_{s_n}} \sum_{j=1}^J w_j \|\mathbf{y} - \mathcal{A}(\hat{\mathbf{X}}_1^{(j)})\|_2^2$
- 11:   **end for**
- 12:    $\hat{\mathbf{x}}_{k+1} \leftarrow \mathbf{X}_{s_N}$
- 13: **end for**
- 14: **return**  $\{\hat{\mathbf{x}}_k\}_{k=1}^K$

**Algorithm**



Paper



GitHub

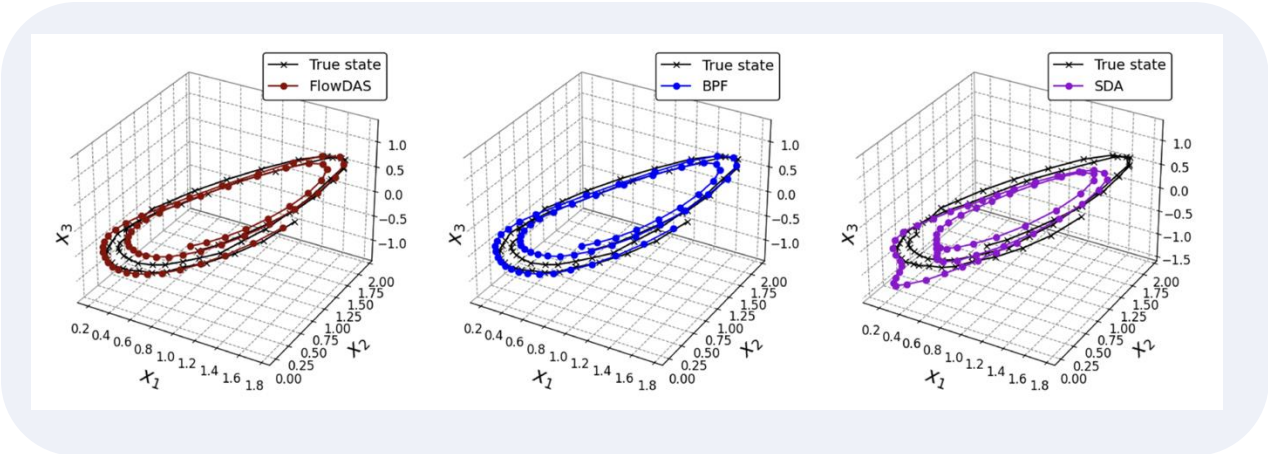


FlowDAS has been accepted by the NeurIPS 2025!

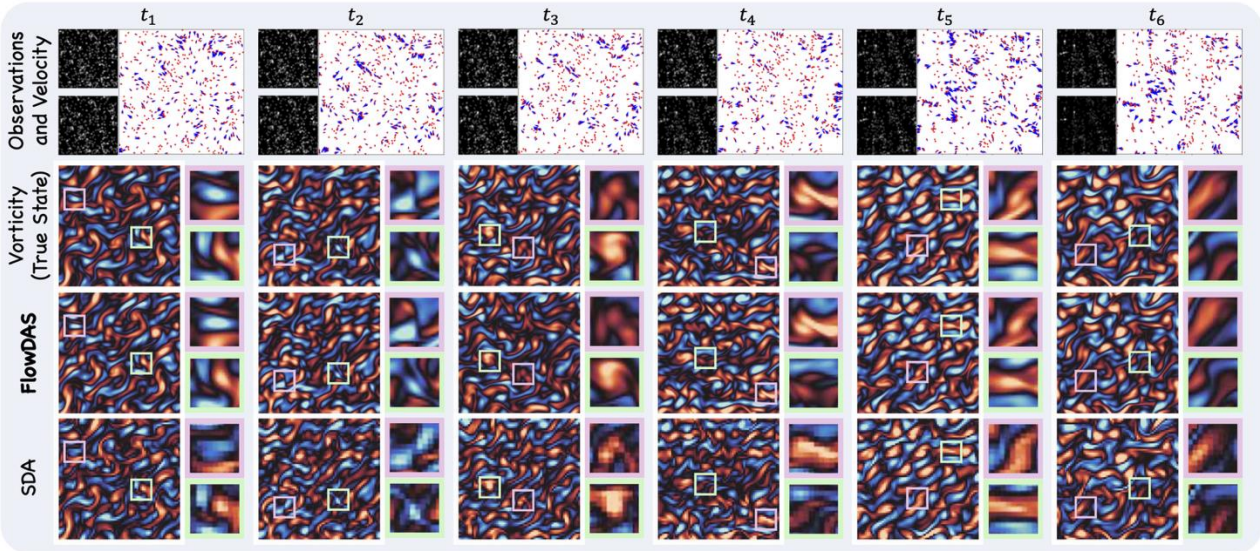




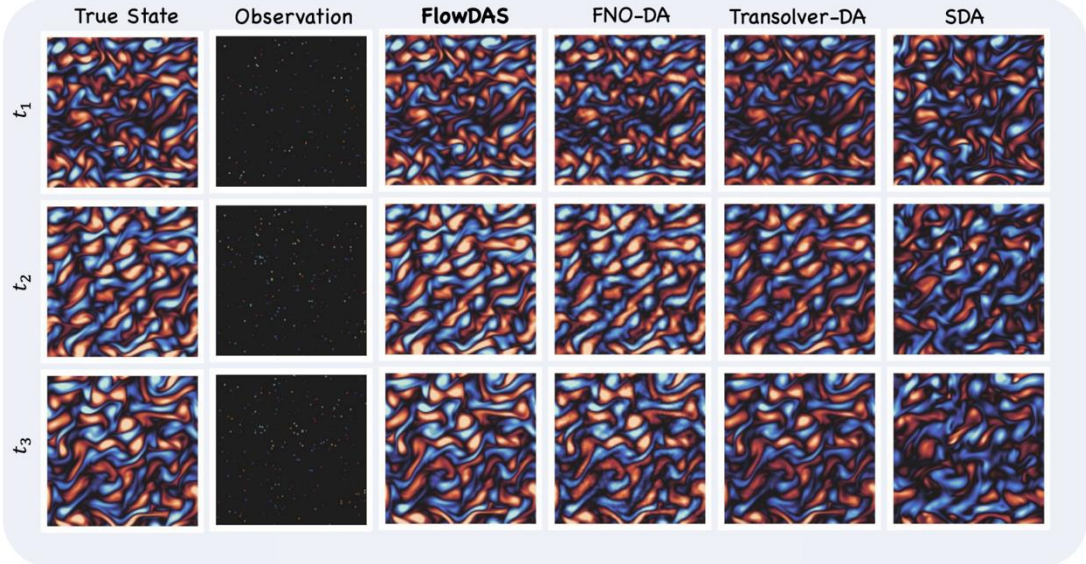
# FlowDAS: A Stochastic Interpolants-based Framework for Data Assimilation



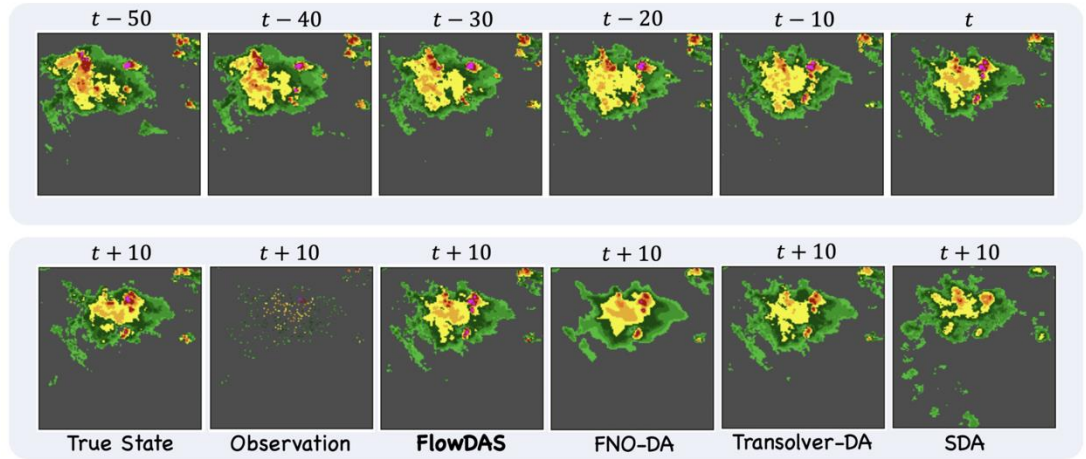
Experiments on **Lorenz-63** data assimilation task



Experiments on **Particle Image Velocimetry**



Experiments on **Navier-Stokes** sparse observation task



Experiments on **SEVIR** weather forecasting task



# Probabilistic Error Analysis of a Randomized Proper Orthogonal Decomposition



• Kathrin Smetana<sup>1</sup>, Tommaso Taddei<sup>2</sup>, Marissa Whitby<sup>1</sup> and Zhiyu Yin<sup>1</sup>

• <sup>1</sup> Stevens Institute of Technology

<sup>2</sup> Inria Bordeaux South-West

**Main goals:** To approximate the range of a bounded nonlinear function  $u : P \rightarrow H$  using randomized POD and derive error bounds that do not rely on the dimension of the ambient space.

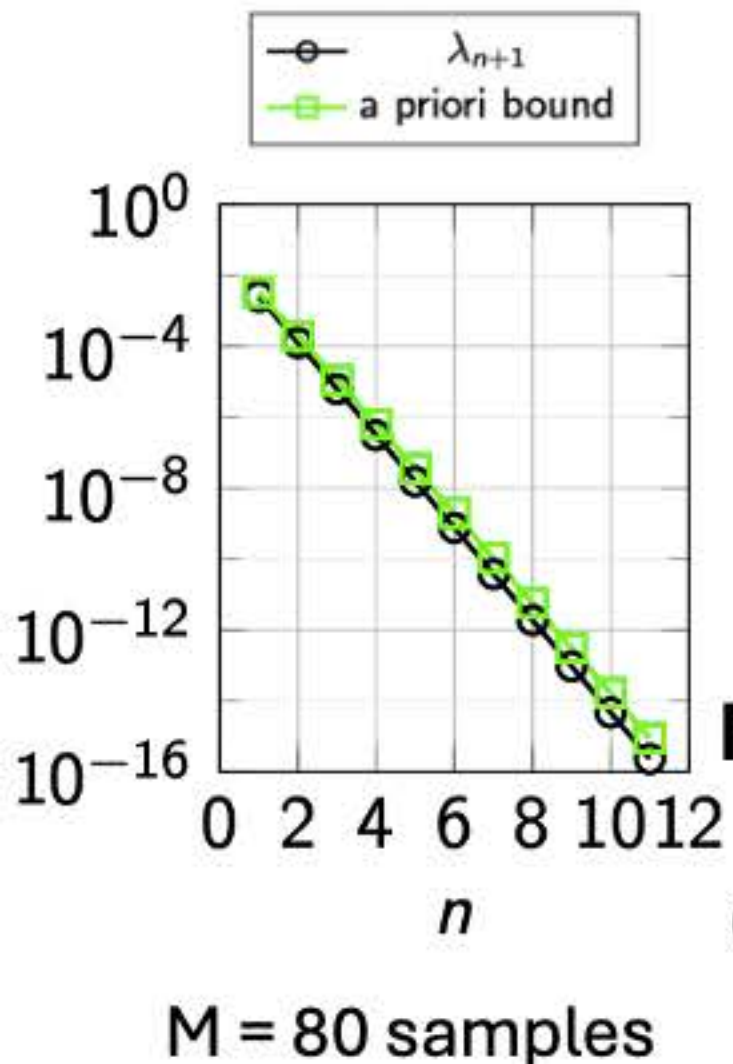
**Current Result:** We have a nonasymptotic probabilistic error bound of the covariance estimator used in randomized POD/PCA for bounded random variables.

**Current Direction:** We are currently using arguments from [Reiß, Wahl 20] to derive a priori error bounds for randomized POD/PCA for bounded random variables.

**Motivation:** For a more restrictive subclass of sub-Gaussian random variables, which excludes some bounded random variables, [Reiß, Wahl 20] shows if a covariance operator has exponentially decaying eigenvalues, then the number of samples needed for POD/PCA to produce (quasi-)optimal results scales linearly in the reduced space.

## Applications:

- PCA of gradient of log-likelihood function for Bayesian Inverse Problems [Zahm, Cui, Law, Spantini, Marzouk 22]
- Construction of local multiscale ansatz spaces [Smetana, Taddei 23]



Supported by NSF Award # 2145364



# The Problem of High-Dimensional Summation

Why Naive Summation Fails for Low-Rank Tucker Tensors

## The Problem: Rank Inflation

Summing low-rank Tucker tensors,  $\mathcal{C} = \sum_{i=1}^d \mathcal{X}^{(i)}$ , is computationally challenging.

- The core issue is **Rank Inflation**.
- The rank of the sum becomes the sum of individual ranks:  $R_k \approx \sum_i r_k^{(i)}$ .
- This causes exponential growth in storage and computational cost.

## The Conventional, Inefficient Solution

This leads to a costly process:

- 1 **Sum:** Ranks additively combine, creating a high-rank tensor.
- 2 **Round:** An expensive decomposition is required to compress the result.

**The question:** Can we find the final low-rank sum *without* forming the high-rank intermediates?

# Breaking the Cycle with Structured Sketching

A Shortcut to the Sum

## Proposed Solution: Sketch, then Sum

Our approach bypasses the need for rounding.

- We create small, compressed **sketches** of each tensor's unfolding:  $\mathbf{Y}_k^{(i)} = \mathbf{X}_{(k)}^{(i)} \mathbf{S}_k$ .
- We sum these small sketches to find an approximate basis,  $\{\mathbf{U}_k\}$ , for the final sum.
- **Reconstruction:** We project each original core tensor onto this new basis and sum the results to form the final core  $\tilde{\mathcal{G}}$ .
- **The result:** An accurate, low-rank approximation  $\tilde{\mathcal{C}} = \tilde{\mathcal{G}} \times_1 \mathbf{U}_1 \cdots \times_N \mathbf{U}_N$  with none of the expensive detours.

## The Theoretical Core & Next Frontier

**The Key Insight:** Method efficacy hinges on choosing the sketch size. Our **Sub-rank Selection** heuristic provides a principled way to balance accuracy and cost.

### Interesting discussions:

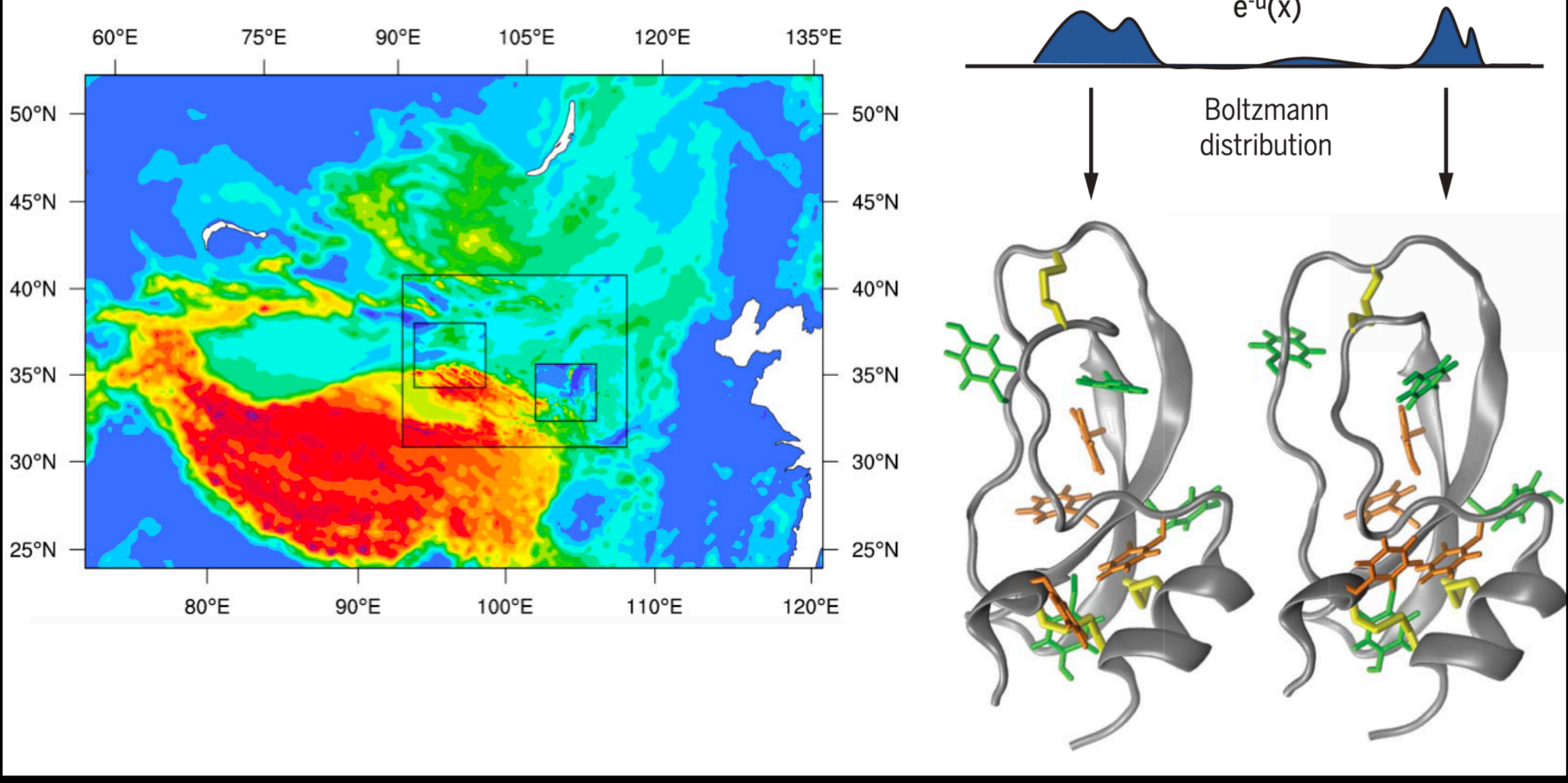
- How do we find the optimal sketch size?
- Can this process be made fully adaptive?
- What are the next steps for scalable tensor computations; a hierarchical Approach?

# Identification of Paths for Dynamic Measure Transport

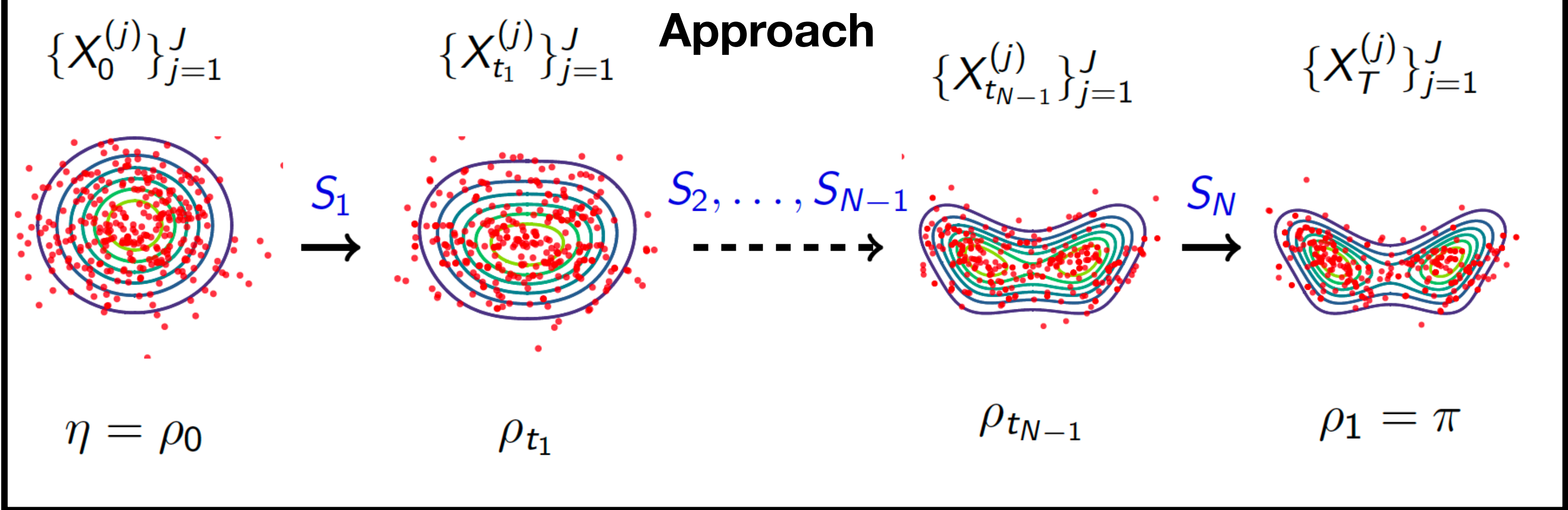
Aimee Maurais, Bamdad Hosseini, and Youssef Marzouk

Task: Sample from a distribution  $\pi$  on  $\mathbb{R}^d$

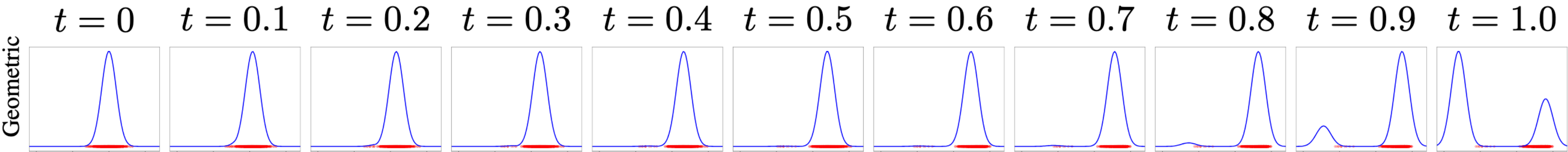
## Motivation



## Approach



The commonly used geometric mixture  $\rho(\cdot, t) \propto \eta^{1-t} \pi^t$  may be problematic!



Can we identify a path  $\rho$  for which a corresponding velocity field is “nice”?

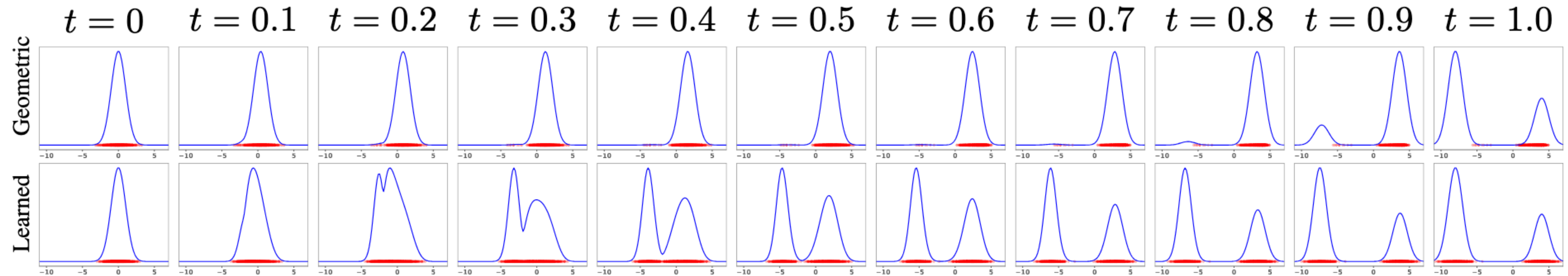
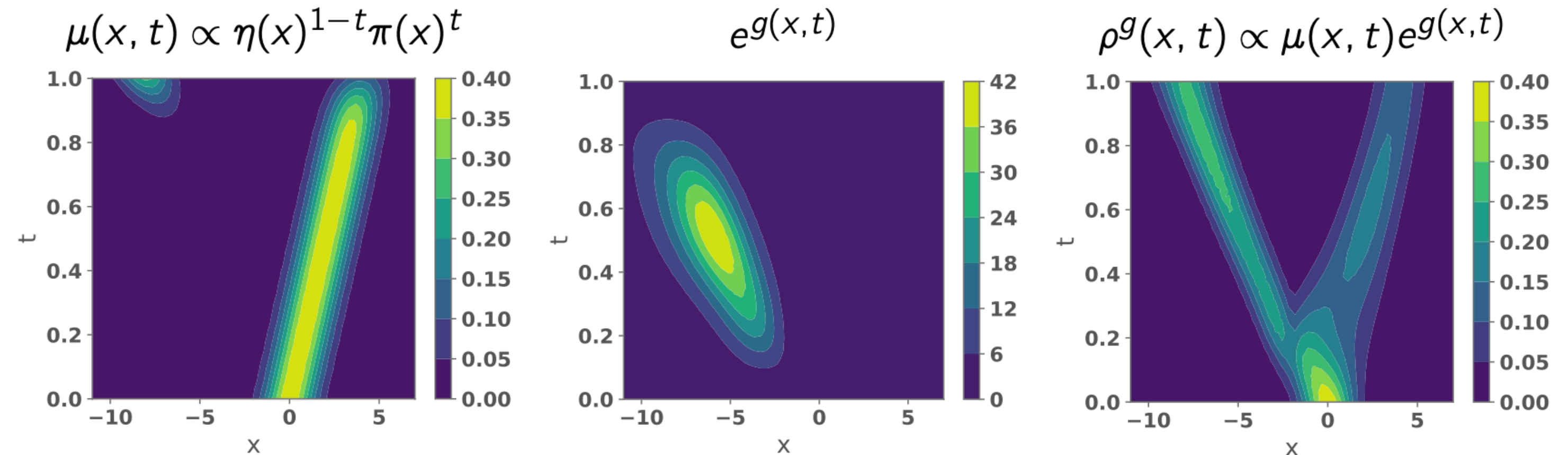


# Approach: Identify a *tilting* via control

We seek a *new* path  $\rho^g \propto \rho_{\text{ref}} e^g$ , where  $\rho_{\text{ref}}$  is a reference path, by solving

$$\inf_{v \in \mathcal{V}, g \in \mathcal{G}} \|v\|_{\mathcal{V}}^2 + \lambda_g \|g\|_{\mathcal{G}}^2 \quad \text{s.t.} \quad -\nabla \cdot (v \rho^g) = \rho^g (\partial_t \log \rho^g), \quad \rho^g \propto \rho^{\text{ref}} e^g, \quad g(\cdot, 0) = g(\cdot, 1) \equiv 0$$

- Formalizes previous approaches to fix the geometric mixture
- Flexible framework enables promotion of *smoothness*
- Many implementations possible!





# Bayesian Inference for Latent Gaussian Models Governed by PDEs

Sonia Reilly and Georg Stadler  
Courant Institute, NYU

Hierarchical model with Gaussian prior:

$$\begin{aligned}\boldsymbol{\theta} &\sim \pi_{\text{hyp}}(\boldsymbol{\theta}) \\ \boldsymbol{m}|\boldsymbol{\theta} &\sim \mathcal{N}(\boldsymbol{\mu}_{\text{pr}}(\boldsymbol{\theta}), \boldsymbol{Q}_{\text{pr}}^{-1}(\boldsymbol{\theta})) \\ \boldsymbol{y}|\boldsymbol{m}, \boldsymbol{\theta} &\sim \pi_{\text{like}}(\boldsymbol{y}|\boldsymbol{m}, \boldsymbol{\theta})\end{aligned}$$

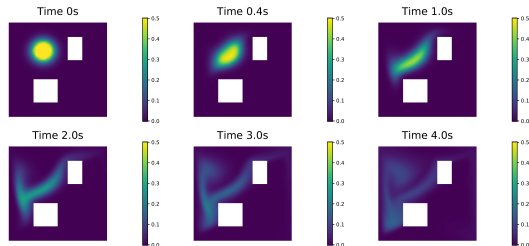
Linear Gaussian Bayesian inverse problem as LGM:

$$\boldsymbol{y} = \boldsymbol{A}\boldsymbol{m} + \boldsymbol{\varepsilon} \quad \text{with} \quad \boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \boldsymbol{Q}_{\boldsymbol{\varepsilon}}(\boldsymbol{\theta})),$$

with  $\boldsymbol{A}$  a discretization of a linear PDE.

Want to characterize  $\pi(\boldsymbol{m}|\boldsymbol{y})$ .

E.g. Advection-Diffusion



Given later time observations  $\boldsymbol{y}$  and prior with unknown hyperparameters  $\boldsymbol{\theta}$ , find posterior of initial condition  $\boldsymbol{m}$

## Fast Sampling Using Hyperparameter Marginal $\pi(\boldsymbol{\theta}|\mathbf{y})$

Sample  $\mathbf{m}^* \sim \pi(\mathbf{m}|\mathbf{y})$  by

1. sampling  $\boldsymbol{\theta}^* \sim \pi(\boldsymbol{\theta}|\mathbf{y})$  (low-dimensional, so can use MCMC)
2. sampling  $\mathbf{m}^* \sim \pi(\mathbf{m}|\boldsymbol{\theta}^*, \mathbf{y})$  (Gaussian, since it is the posterior of a linear Gaussian Bayesian inverse problem)

Need to compute quickly:

$$\begin{aligned}\pi(\boldsymbol{\theta}|\mathbf{y}) &\propto \frac{\pi(\mathbf{m}, \boldsymbol{\theta}, \mathbf{y})}{\pi(\mathbf{m}|\boldsymbol{\theta}, \mathbf{y})} = \frac{\pi_{\text{like}}(\mathbf{y}|\mathbf{m}, \boldsymbol{\theta})\pi_{\text{pr}}(\mathbf{m}|\boldsymbol{\theta})\pi_{\text{hyp}}(\boldsymbol{\theta})}{\pi(\mathbf{m}|\boldsymbol{\theta}, \mathbf{y})} \\ &\propto \left( \frac{|\mathbf{Q}_{\text{pr}}||\mathbf{Q}_{\varepsilon}|}{|\mathbf{Q}_{\text{post}}|} \right)^{1/2} \exp \left( -\frac{1}{2} \left[ \|\mathbf{y}\|_{\mathbf{Q}_{\varepsilon}} + \|\boldsymbol{\mu}_{\text{pr}}\|_{\mathbf{Q}_{\text{pr}}} - \|\boldsymbol{\mu}_{\text{post}}\|_{\mathbf{Q}_{\text{post}}} \right] \right) \pi_{\text{hyp}}(\boldsymbol{\theta})\end{aligned}$$

Idea: low rank approximation of  $\mathbf{Q}_{\text{post}} - \mathbf{Q}_{\text{pr}} = \mathbf{A}^T \mathbf{Q}_{\varepsilon} \mathbf{A}$  (two ways)