

Training-Free Bayesian Filtering with Generative Emulators

*Forty-third International Conference on Machine Learning
Poster Session 5 - Wednesday 8 July 5 p.m. - Poster #1405*



Thomas Savary¹



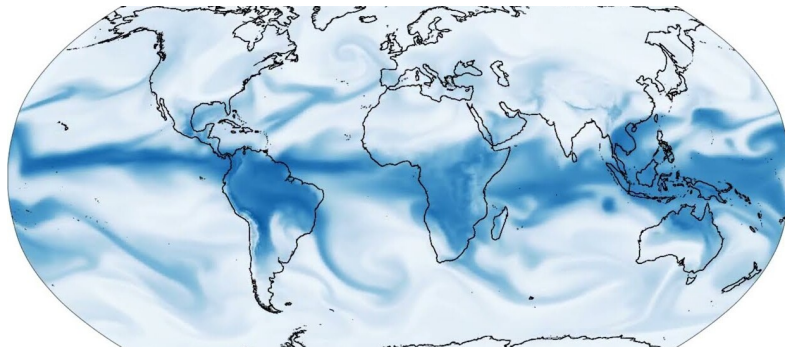
François Rozet¹



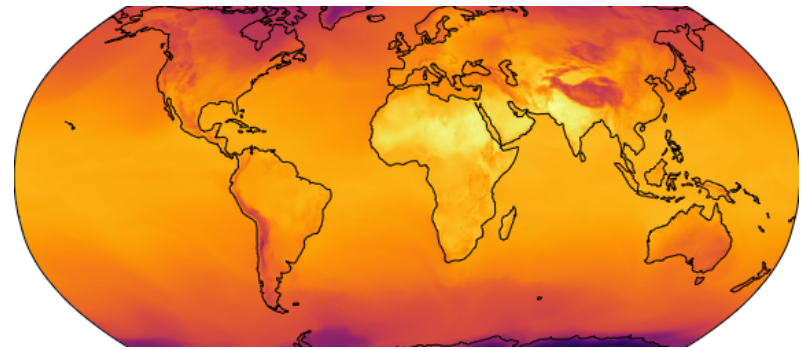
Gilles Louppe¹

¹SAIL, Montefiore Institute, University of Liège, Belgium

Over the past decades, weather forecasting has improved dramatically, and forecasts can now be accessed almost instantly from our smartphones.



Specific humidity of the atmosphere at 700 hPa



Surface temperature

These advances have been made possible by continuous **improvements in dynamical models** and **advances in computing capabilities**.

The atmosphere is discretized on a grid, and, **starting from an initial condition**, we solve partial differential equations to simulate its future evolution.



The atmosphere is discretized on a grid, and, **starting from an initial condition**, we solve partial differential equations to simulate its future evolution.



! Estimating this initial condition is challenging.

The atmosphere is discretized on a grid, and, **starting from an initial condition**, we solve partial differential equations to simulate its future evolution.



$$\frac{\partial u}{\partial t} + (u \cdot \nabla)u = -\frac{\nabla p}{\rho} + \nu \nabla^2 u$$

Prior knowledge of the system dynamics

! Estimating this initial condition is challenging. Bayesian filtering addresses this problem by combining **prior knowledge of the dynamics**

The atmosphere is discretized on a grid, and, **starting from an initial condition**, we solve partial differential equations to simulate its future evolution.



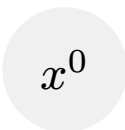
$$\frac{\partial u}{\partial t} + (u \cdot \nabla)u = -\frac{\nabla p}{\rho} + \nu \nabla^2 u$$

Prior knowledge of the system dynamics

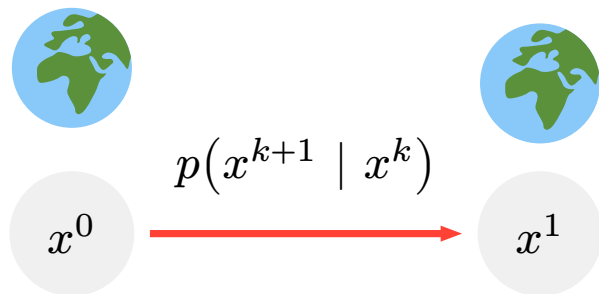


Noisy and sparse observations

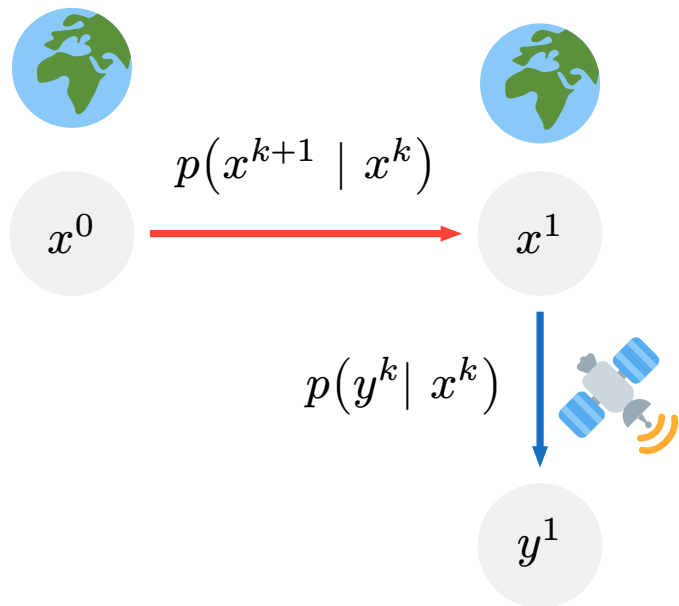
! Estimating this initial condition is challenging. Bayesian filtering addresses this problem by combining **prior knowledge of the dynamics** with **observations**.



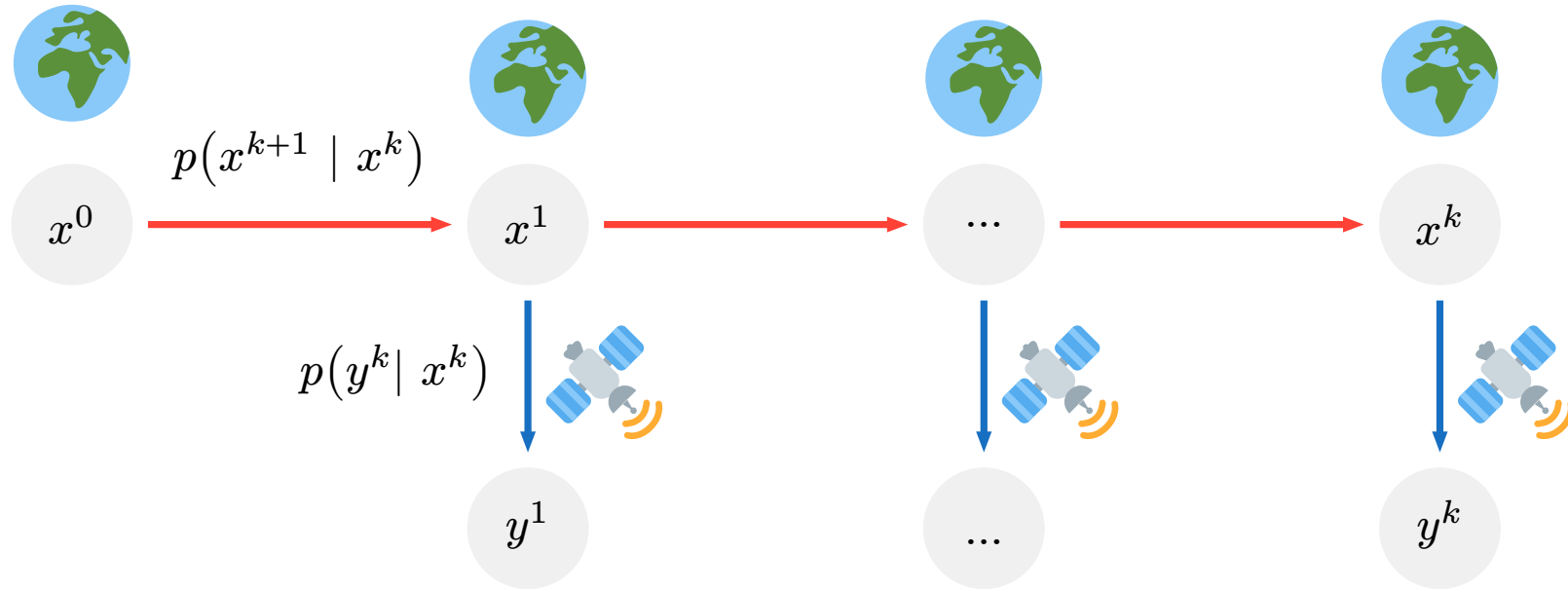
Given a **prior distribution** $p(x_0)$,



Given a **prior distribution** $p(x_0)$, a **Markovian transition model** $p(x^{k+1} | x^k)$

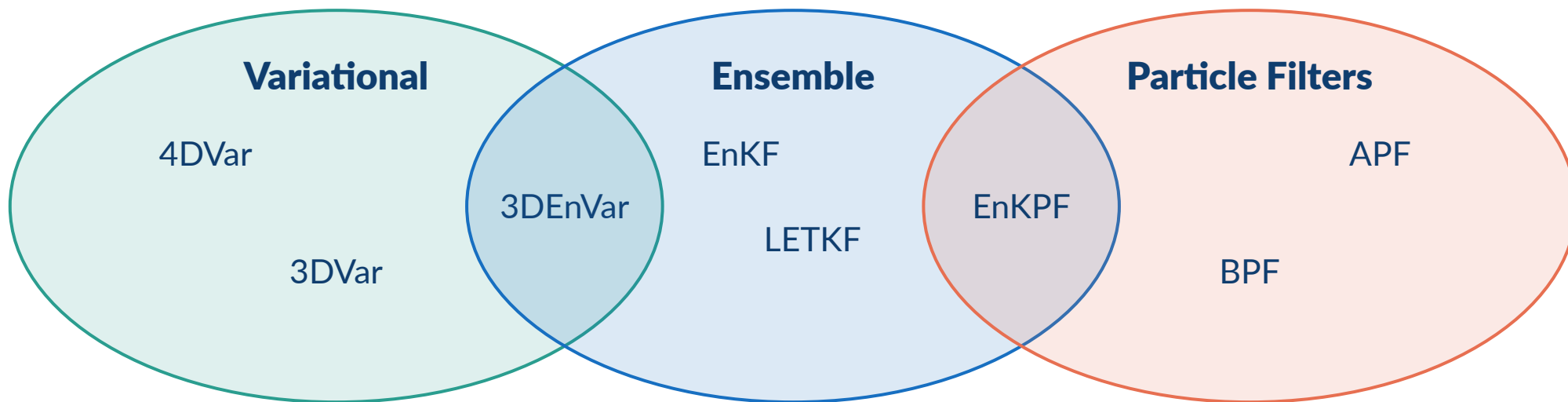


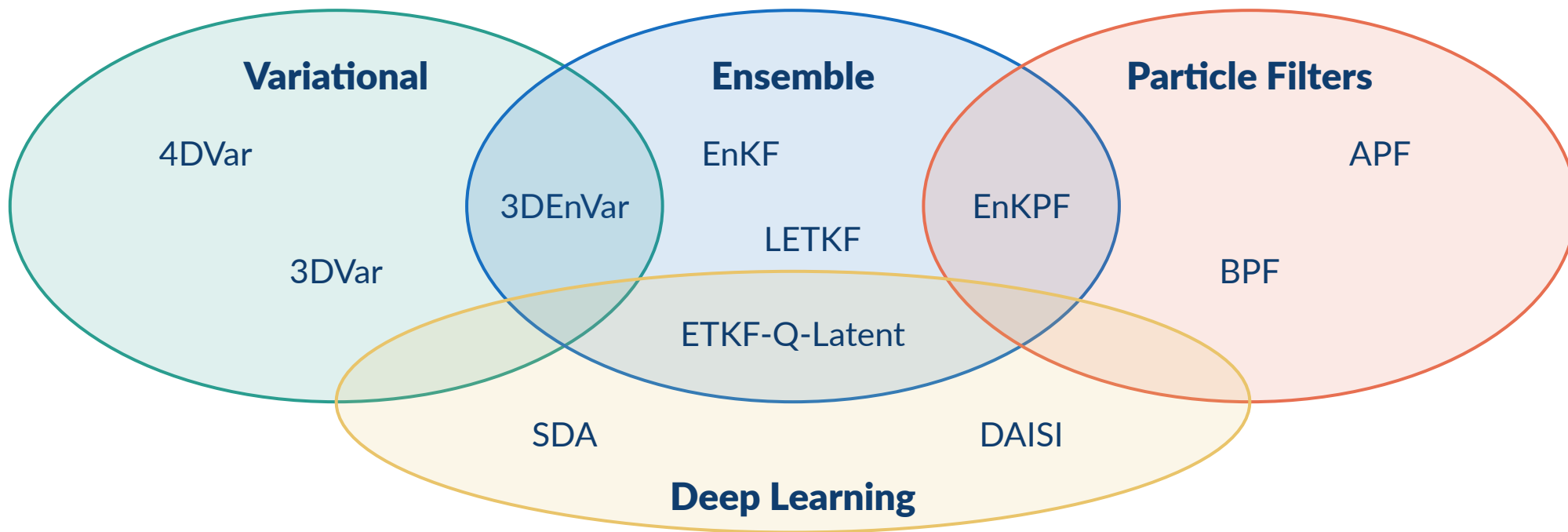
Given a **prior distribution** $p(x_0)$, a **Markovian transition model** $p(x^{k+1} | x^k)$ and an **observation model** $p(y^k | x^k)$,

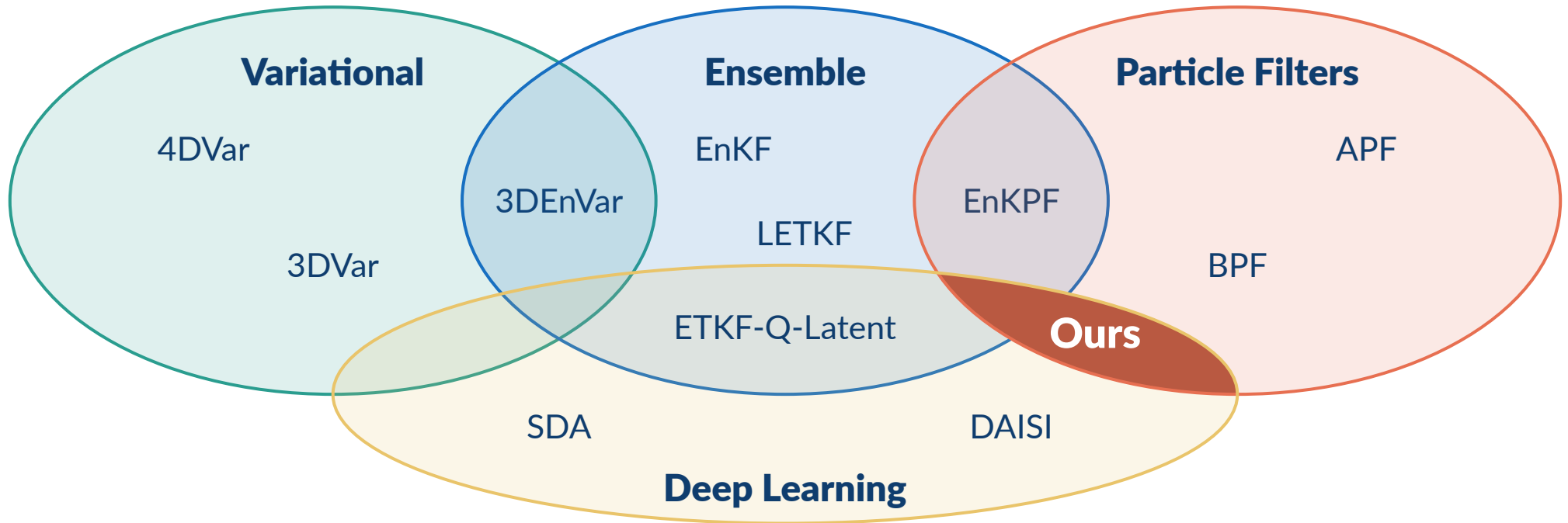


Given a **prior distribution** $p(x_0)$, a **Markovian transition model** $p(x^{k+1} | x^k)$ and an **observation model** $p(y^k | x^k)$, the goal of **Bayesian filtering** is to estimate plausible states from current and past observations

$$p(x^k | y^{1:k}) = p(x^k | y^1, \dots, y^k) \propto p(y^k | x^k) \int p(x^k | x^{k-1}) p(x^{k-1} | y^{1:k-1}) dx^{k-1}.$$

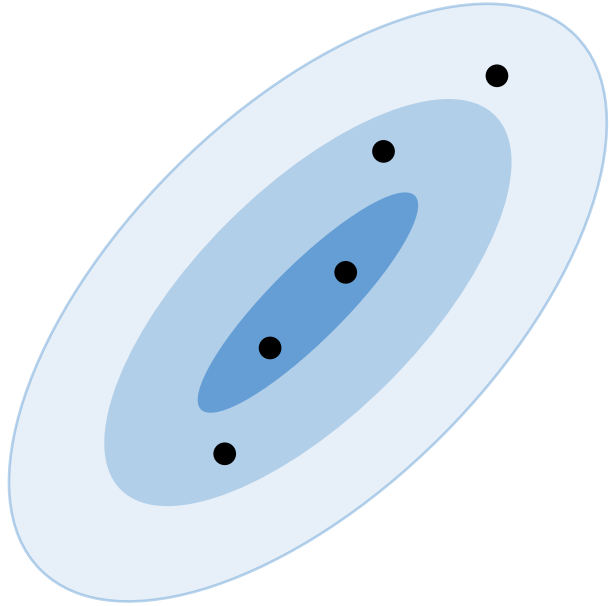






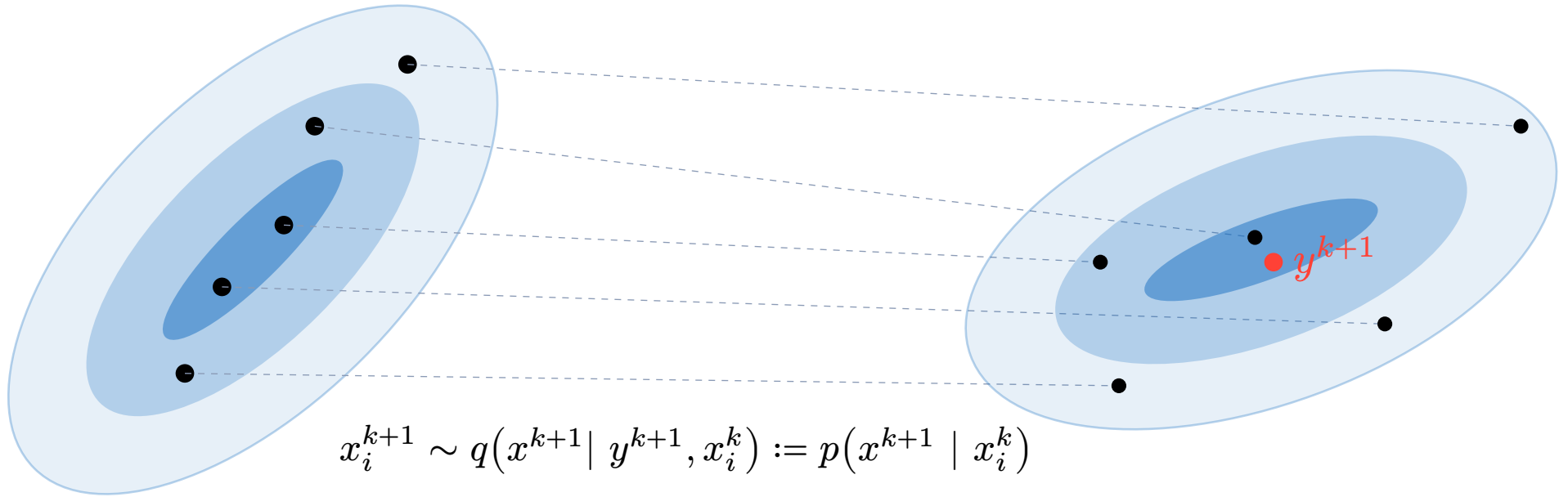
“Finally, the community anticipates that ML will help address longstanding limitations of traditional DA [...]. For example, ML may help operationalize fully non-Gaussian DA methods such as particle filters, which have historically been limited by computational cost in high-dimensional systems.”

Particle filters are iterative algorithms that approximate the posterior distribution $p(x^k \mid y^{1:k})$ by a discrete probability measure.



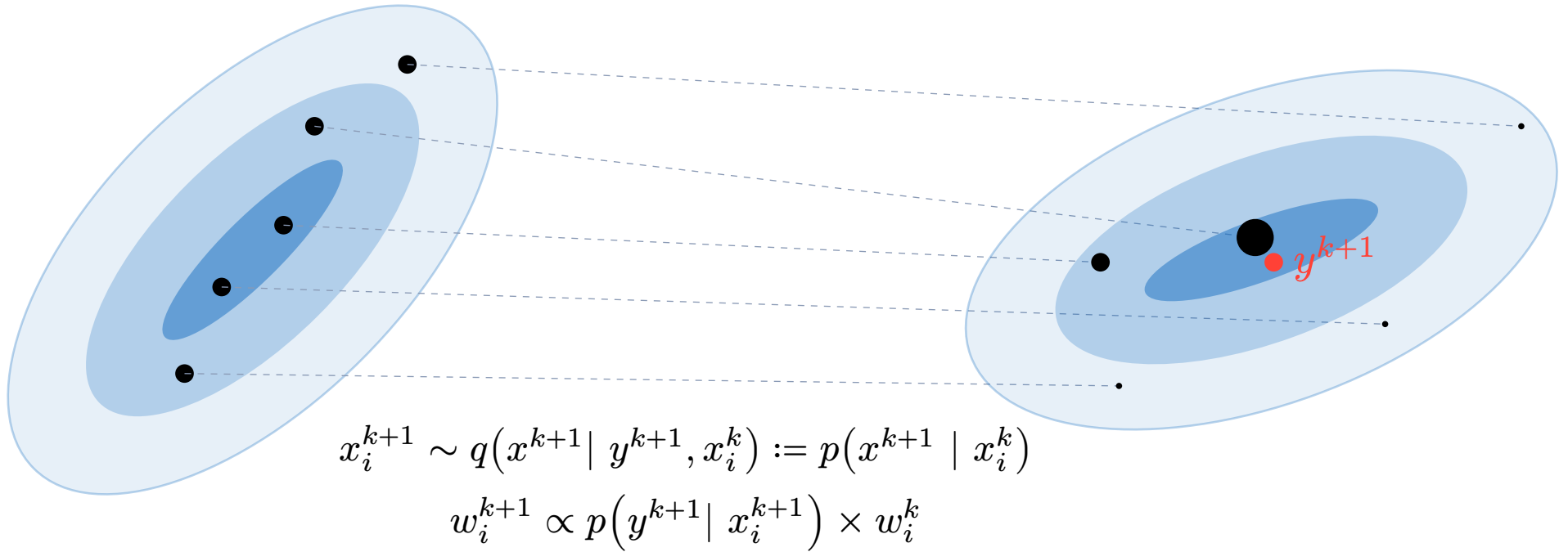
$$p(x^k \mid y^{1:k}) \approx \sum_{i=1}^N w_i^k x_i^k$$

Particle filters are iterative algorithms that approximate the posterior distribution $p(x^k \mid y^{1:k})$ by a discrete probability measure. Given a new **observation** y^{k+1} , particles are propagated using a **proposal distribution** $q(x^{k+1} \mid y^{k+1}, x^k)$



$$p(x^k \mid y^{1:k}) \approx \sum_{i=1}^N w_i^k x_i^k$$

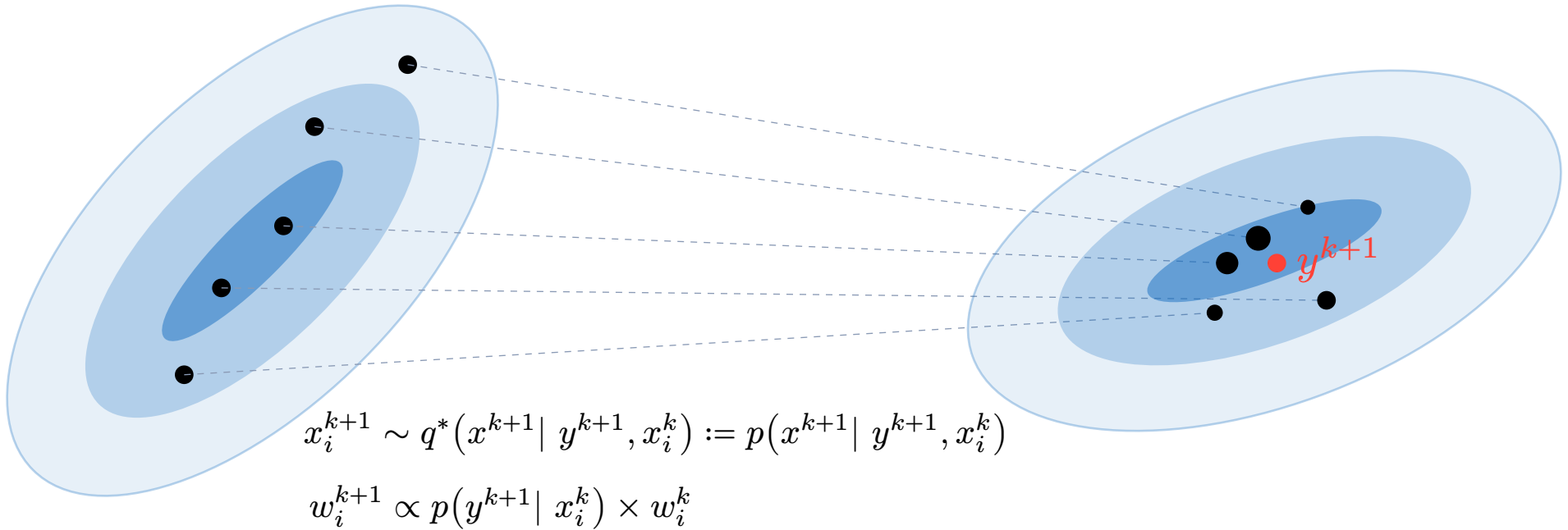
Particle filters are iterative algorithms that approximate the posterior distribution $p(x^k \mid y^{1:k})$ by a discrete probability measure. Given a new **observation** y^{k+1} , particles are propagated using a **proposal distribution** $q(x^{k+1} \mid y^{k+1}, x^k)$ and **weighted** to correct the mismatch.



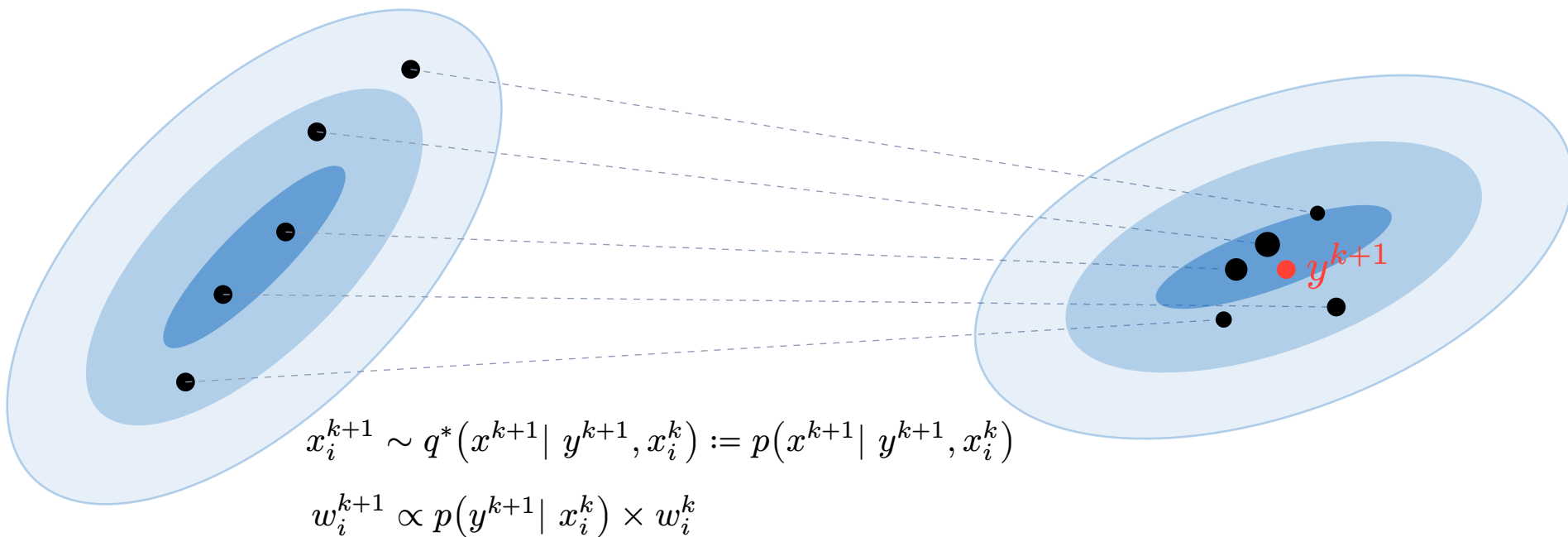
$$p(x^k \mid y^{1:k}) \approx \sum_{i=1}^N w_i^k x_i^k$$

$$p(x^{k+1} \mid y^{1:k+1}) \approx \sum_{i=1}^N w_i^{k+1} x_i^{k+1}$$

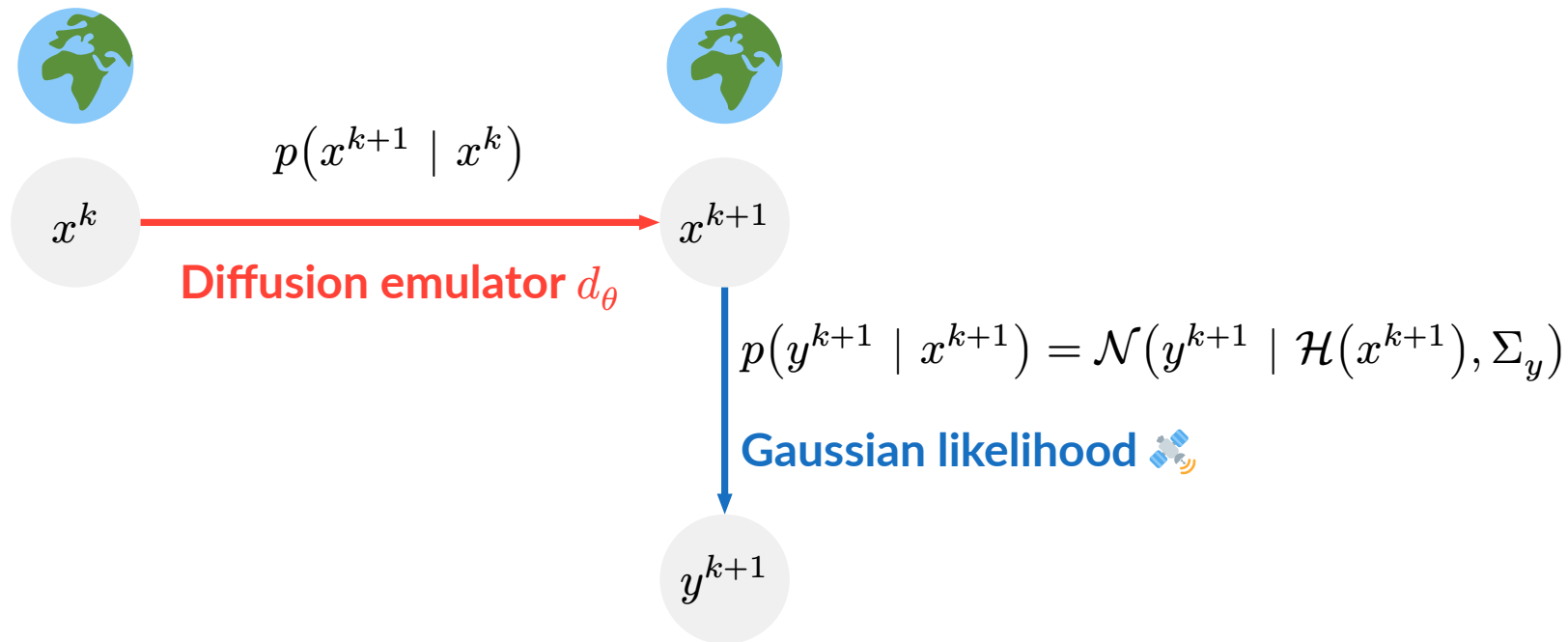
It has been shown in the particle filters literature (Snyder et al, 2015) that the **optimal proposal** $q^*(x^{k+1} | y^{k+1}, x^k) = p(x^{k+1} | y^{k+1}, x^k)$ minimizes weights degeneracy.



It has been shown in the particle filters literature (Snyder et al, 2015) that the **optimal proposal** $q^*(x^{k+1} | y^{k+1}, x^k) = p(x^{k+1} | y^{k+1}, x^k)$ minimizes weights degeneracy.



⚠ It is (very) difficult to sample from the optimal distribution and compute the associated weights for classical numerical solvers, but (almost) not for diffusion emulators.



💡 Given a **diffusion emulator** for the transition $p(x^{k+1} | x^k)$, we can perform approximate **zero-shot sampling from the optimal proposal** $p(x^{k+1} | y^{k+1}, x^k)$ (i.e. without further training).

Starting from a noisy sample $x_{t=1}^{k+1}$ and given a particle x_i^k , we can draw sample from the optimal proposal $p(x^{k+1} \mid y^{k+1}, x_i^k)$ by solving the following SDE from $t = 1$ to $t = 0$

$$dx_t^{k+1} = \left[f_t x_t^{k+1} - g_t^2 \nabla_{x_t^{k+1}} \log p_t(x_t^{k+1} \mid y^{k+1}, x_i^k) \right] dt + g_t dw_t.$$

Starting from a noisy sample $x_{t=1}^{k+1}$ and given a particle x_i^k , we can draw sample from the optimal proposal $p(x^{k+1} \mid y^{k+1}, x_i^k)$ by solving the following SDE from $t = 1$ to $t = 0$

$$dx_t^{k+1} = \left[f_t x_t^{k+1} - g_t^2 \nabla_{x_t^{k+1}} \log p_t(x_t^{k+1} \mid y^{k+1}, x_i^k) \right] dt + g_t dw_t.$$

The posterior score $\nabla_{x_t^{k+1}} \log p_t(x_t^{k+1} \mid y^{k+1}, x_i^k)$ is unknown a priori, but can be decomposed via Bayes' rule as

$$\nabla_{x_t^{k+1}} \log p_t(x_t^{k+1} \mid x_i^k) + \nabla_{x_t^{k+1}} \log p_t(y^{k+1} \mid x_t^{k+1}, x_i^k).$$

Starting from a noisy sample $x_{t=1}^{k+1}$ and given a particle x_i^k , we can draw sample from the optimal proposal $p(x^{k+1} \mid y^{k+1}, x_i^k)$ by solving the following SDE from $t = 1$ to $t = 0$

$$dx_t^{k+1} = \left[f_t x_t^{k+1} - g_t^2 \nabla_{x_t^{k+1}} \log p_t(x_t^{k+1} \mid y^{k+1}, x_i^k) \right] dt + g_t dw_t.$$

The posterior score $\nabla_{x_t^{k+1}} \log p_t(x_t^{k+1} \mid y^{k+1}, x_i^k)$ is unknown a priori, but can be decomposed via Bayes' rule as

$$\nabla_{x_t^{k+1}} \log p_t(x_t^{k+1} \mid x_i^k) + \nabla_{x_t^{k+1}} \log p_t(y^{k+1} \mid x_t^{k+1}, x_i^k).$$

The first term is already known using the trained denoiser, and following Rozet et al. (NeurIPS 2024), we approximate the second term by

$$\nabla_{x_t^{k+1}}^T \mathbb{E}[x^{k+1} \mid x_t^{k+1}, x_i^k] H^T (\Sigma_y + H \mathbb{V}[x^{k+1} \mid x_t^{k+1}, x_i^k] H^T)^{-1} (y - H \mathbb{E}[x^{k+1} \mid x_t^{k+1}, x_i^k]).$$

⚠ Computing particle weights for the optimal proposal is not trivial since we need to evaluate

$$w_i^{k+1} \propto p(y^{k+1} \mid x_i^k) = \int p(y^{k+1} \mid x^{k+1}) p(x^{k+1} \mid x_i^k) dx^{k+1}.$$

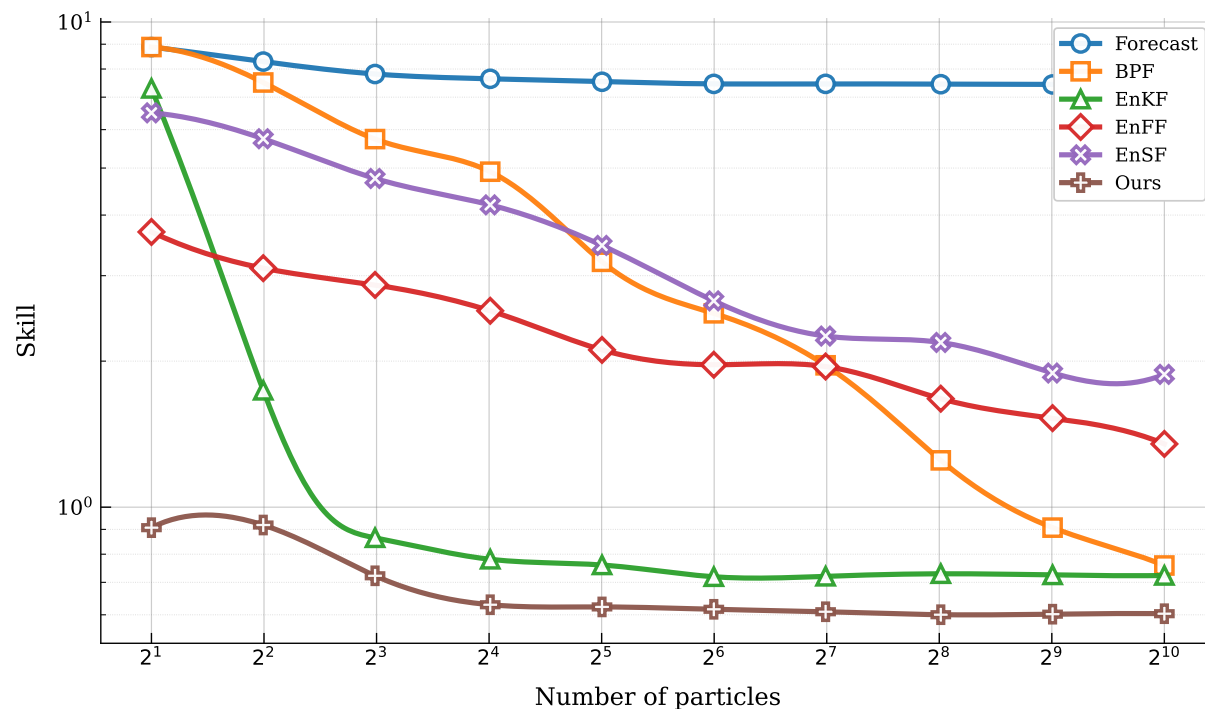
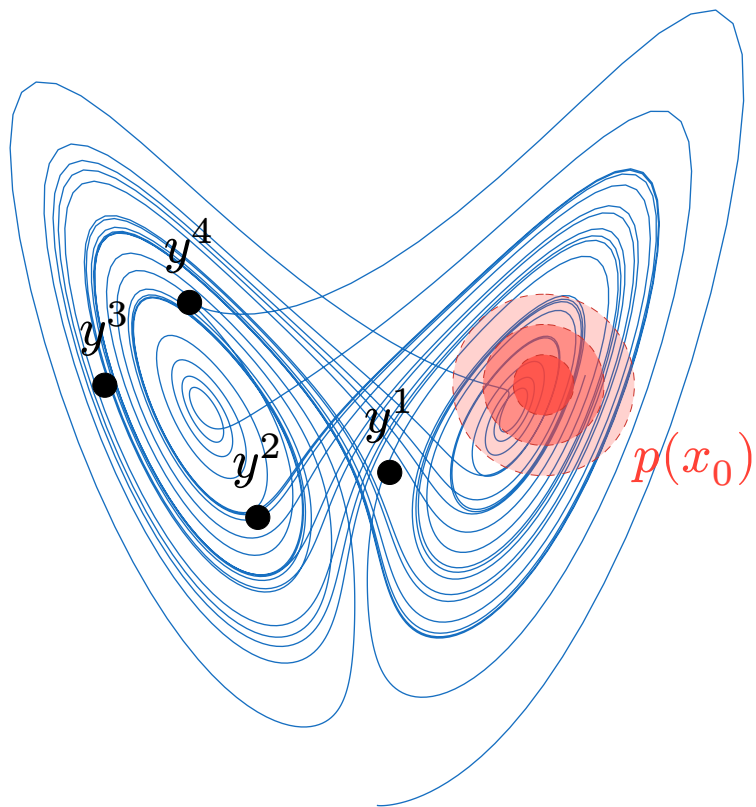
Assuming that the transition distribution $p(x^{k+1} \mid x^k)$ is a Dirac at $\mathbb{E}[x^{k+1} \mid x_i^k]$, we obtain

$$w_i^{k+1} \propto \|y^{k+1} - \mathcal{H}(\mathbb{E}[x^{k+1} \mid x_i^k])\|_{\alpha^{-1}\Sigma_y},$$

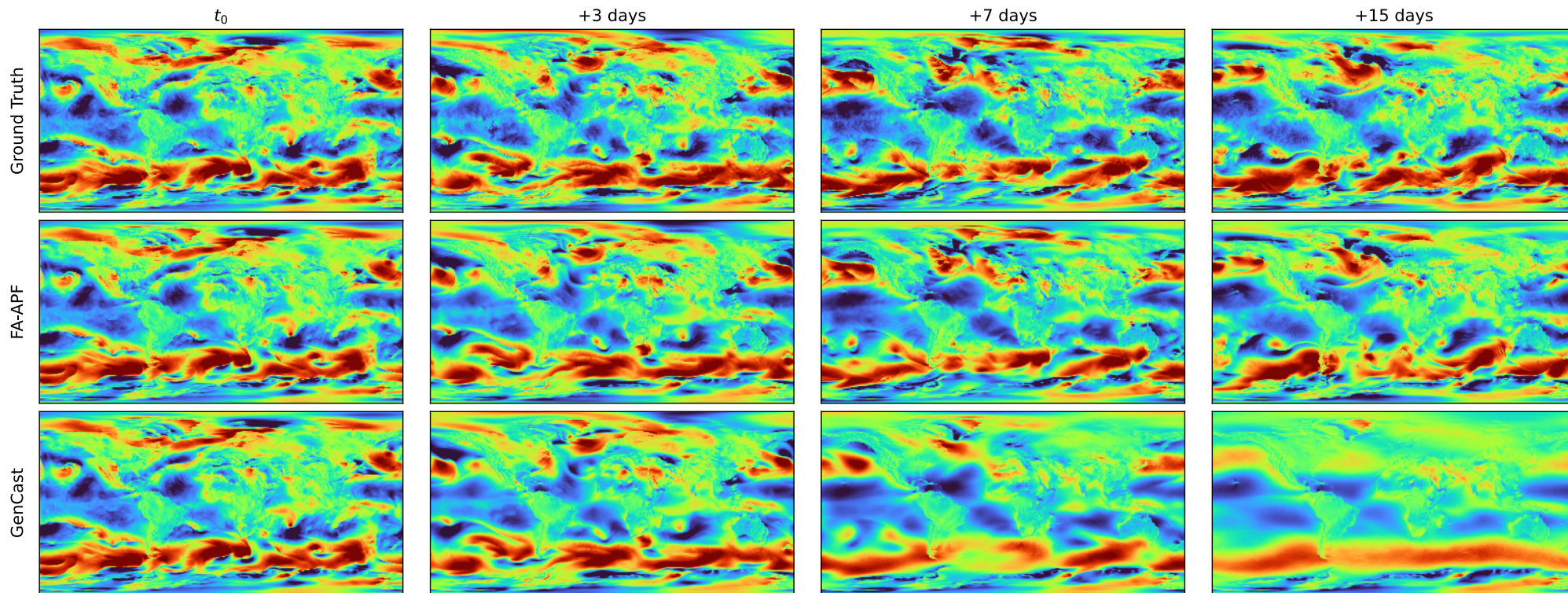
and we can compute the expectation of the next state with only one call to the denoiser

$$\mathbb{E}[x^{k+1} \mid x_i^k] \stackrel{\varepsilon \sim \mathcal{N}(0, I)}{=} \mathbb{E}[x^{k+1} \mid x_i^k, \sigma_1 \varepsilon] \approx d_\theta(t=1, x^k = x_i^k, x_t^{k+1} = \sigma_1 \varepsilon).$$

We first evaluate our approach on a stochastic version of the **Lorenz-63** system (Chekroun et al, 2011). **Our method is more effective than the BPF**, especially with a small number of particles.



We also applied our method to the atmosphere using the pre-trained GenCast (Price et al, 2023) denoiser, enabling the assimilation of sparse observations inspired by real stations.



Key Takeaways

- Particle filters are **exact in nonlinear regimes**, but suffer from **degeneracy**.
- **Diffusion emulators** enable sampling from the **optimal proposal** without further training.
- Experiments on large-scale dynamical systems demonstrate the potential of this method.

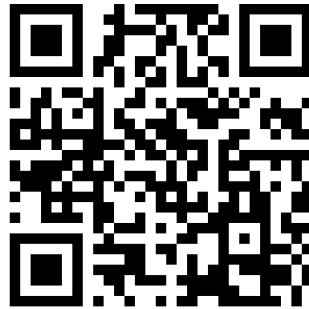
Limitations

- In theory, **particle filters still require a large number of particles** in large-scale systems.
- **Weights are approximated**, and the observation covariance Σ_y is **artificially inflated**.
- Zero-shot posterior sampling with diffusion emulators **is not exact**.

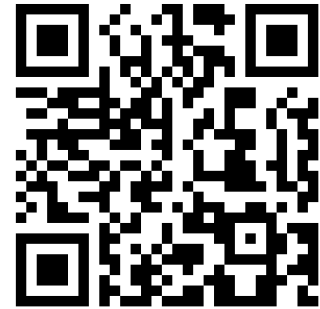
Thank you for your attention!



Paper



Code



LinkedIn

See you at the poster session 😊

Poster #1405