

Learning surrogate models and data assimilation processes for advanced geophysical dynamics forecasting

Marc Bocquet,^{*}

Alban Farchi,^{†,*} Tobias Finn,^{*} Charlotte Durand,^{*} Sibo Cheng,^{*} Wenbo Yu,^{*}
Yumeng Chen,^{||} Ivo Pasmans,^{||} Alberto Carrassi,[§]

^{*} CEREa, ENPC, EDF R&D, Institut Polytechnique de Paris, Île-de-France, France

^{||} Department of Meteorology and National Centre for Earth Observation, University of Reading, United-Kingdom

[§] Department of Physics and Astronomy, University of Bologna, Italy

[†] ECMWF, Reading, United Kingdom,

[‡] NERSC, Bergen, Norway,



- 1 Learning data assimilation from artificial intelligence: first results
 - Sequential data assimilation for chaotic dynamics
 - Learning data assimilation
 - Sensitivity analysis and first results
 - Investigation and interpretation
- 2 Learning data assimilation from artificial intelligence: advanced results
 - Scalability and locality
 - Learning the analysis operator from observations only
 - Learning nonlinearity
- 3 Conclusions

Outline

- 1 Learning data assimilation from artificial intelligence: first results
 - Sequential data assimilation for chaotic dynamics
 - Learning data assimilation
 - Sensitivity analysis and first results
 - Investigation and interpretation
- 2 Learning data assimilation from artificial intelligence: advanced results
 - Scalability and locality
 - Learning the analysis operator from observations only
 - Learning nonlinearity
- 3 Conclusions

Sequential data assimilation for chaotic dynamics

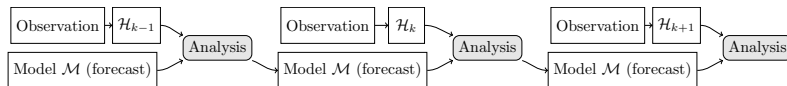
► Here, data assimilation (DA) methods are formulated from

$$\mathbf{x}_{k+1} = \mathcal{M}(\mathbf{x}_k), \quad (1a)$$

$$\mathbf{y}_k = \mathcal{H}_k(\mathbf{x}_k) + \varepsilon_k, \quad \varepsilon_k \sim N(\mathbf{0}, \mathbf{R}_k), \quad (1b)$$

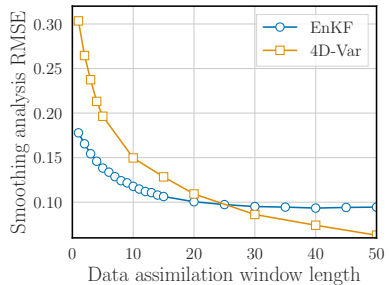
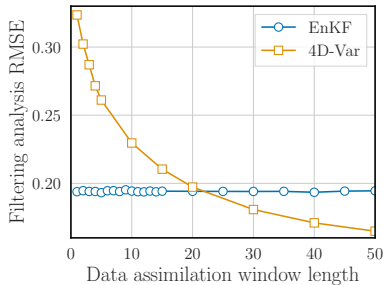
where $\mathcal{M}$ is the *autonomous* evolution model, $\mathbf{x}_k$ is the state vector at time τ_k , $\mathbf{y}_k$ is the observation vector, $\mathcal{H}_k$ is the observation operator, ε_k is the observation error, assumed to be additive, unbiased, white in time, and Gaussian of covariance matrix $\mathbf{R}_k$.

► DA for geofluids has to be *sequential* in time because (i) observations need to be assimilated *as they arrive* to update the state estimation, (ii) applied to *chaotic dynamics*, typical errors have an exponential growth.



The edge of ensemble filtering methods

- **The variational methods (3D-Var, 4D-Var):** can handle nonlinearity of the operators, asynchronous observations, but *cannot handle the errors of the day*.
- **The ensemble filtering methods (EnKFs):** can only handle weak nonlinearity of the operators, cannot handle asynchronous observations, *can handle the errors of the day* through the ensemble.
- Testing the EnKF ($N_e = 20$), 4D-Var, and IEnKS ($N_e = 20$) variants with the chaotic 40-variable Lorenz 96 model [Bocquet et al. 2013; Asch et al. 2016]:



- In mild nonlinear regime, the EnKF significantly outperforms the (basic) 4D-Var with moderately large DA windows because it captures the *errors of the day*.

Our focus: learning the analysis

► Let us assume that $\mathcal{M}$ is known, that the Jacobian of $\mathcal{H}_k$ is $\mathbf{H}_k$, and that we wish to learn an *incremental analysis operator* a_θ , typically a neural network parametrised by θ .

► If $\mathbf{E}_k^a, \mathbf{E}_k^f \in \mathbb{R}^{N_x \times N_e}$ are the analysis and forecast ensemble matrices at time τ_k , a_θ is defined via the (ensemble) update:

$$\mathbf{E}_k^a = \mathbf{E}_k^f + a_\theta \left(\mathbf{E}_k^f, \mathbf{H}_k^\top \mathbf{R}_k^{-1} \delta_k \right), \quad (2a)$$

where δ_k , the innovation at time τ_k , is defined by

$$\delta_k \triangleq \mathbf{y}_k - \mathcal{H}_k \left(\bar{\mathbf{x}}_k^f \right), \quad \bar{\mathbf{x}}_k^f \triangleq \frac{1}{N_e} \sum_{i=1}^{N_e} \mathbf{x}_k^{f,i}. \quad (2b)$$

→ Notice our trick: $a_\theta \left(\mathbf{E}_k^f, \delta_k \right) \longrightarrow a_\theta \left(\mathbf{E}_k^f, \mathbf{H}_k^\top \mathbf{R}_k^{-1} \delta_k \right)$, i.e., *uplift of observation information in state space*.

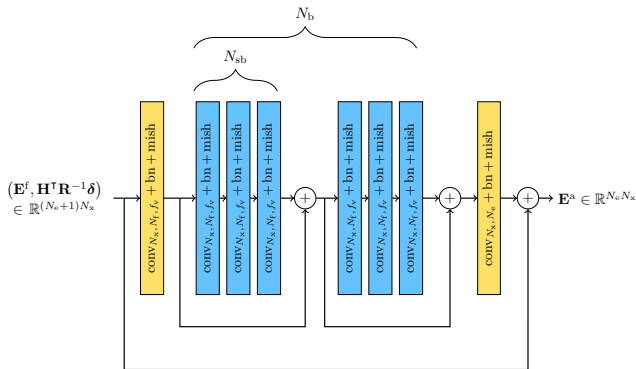
► The DA forecast step propagates the analysis ensemble, member-wise:

$$\mathbf{E}_{k+1}^f = \mathcal{M} \left(\mathbf{E}_k^a \right). \quad (3)$$

► The a_θ -based sequential DA will be called *DAN* in the following.

Neural network architecture

- We choose a_θ to have a simple residual convolutional neural network (CNN) architecture.



Architecture of the residual convolutional network, where $N_b = 2$, $N_{sb} = 3$. $\text{conv}_{N_1, N_2, f}$ is a generic one-dimensional convolutional layer of dimension N_1 , with N_2 filters of kernel size f .

Training scheme – 1/2

- ▶ Full end-to-end, models and DA [Allen et al. 2025; Alexe et al. 2024; Lean et al. 2025]: not our objective here!
- ▶ Literature (focused on *sequential data assimilation*):
 - ▶ Learning the analysis of sequential DA is not new [Härter et al. 2012; Cintra et al. 2018], though barely explored.
 - ▶ Learning key components of the analysis in the (En)KF [H. Hoang et al. 1994; S. Hoang et al. 1998] possibly leveraging auto-differentiable structure [Haarnoja et al. 2016; Chen et al. 2022; Luk et al. 2024] was also investigated.
 - ▶ Only two key papers so far focused on a *non-parametrised* analysis using backpropagation *through the DA cycles*: [McCabe et al. 2021; Boudier et al. 2023].
- ▶ Our training loss (supervised learning):¹

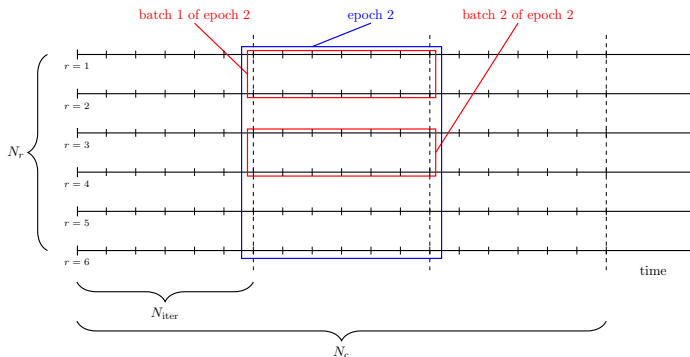
We consider N_r , N_c cycle-long ensemble DA runs, based on N_r independent concurrent trajectories of the dynamics $\mathbf{x}_k^{t,r}$ and as many sequences of observation vectors $\mathbf{y}_k^r$.

The analysis ensemble is $\mathbf{x}_k^{a,i,r} \in \mathbb{R}^{N_x \times N_e}$. The *loss function* is defined by

$$\mathcal{L}(\theta) = \sum_{r=1}^{N_r} \sum_{k=1}^{N_c} \left\| \mathbf{x}_k^{t,r} - \bar{\mathbf{x}}_k^{a,r}(\theta) \right\|^2, \quad \bar{\mathbf{x}}_k^{a,r} \triangleq \frac{1}{N_e} \sum_{i=1}^{N_e} \mathbf{x}_k^{a,i,r}. \quad (4)$$

¹[Bocquet et al. 2024]

Training scheme – 2/2



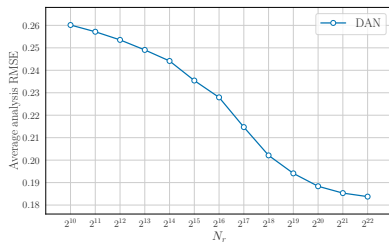
Structure of the dataset organised as a function of time, trajectory sample, batches and epochs.

► Like [McCabe et al. 2021; Boudier et al. 2023], we use *truncated backpropagation through time TBPTT* [Tang et al. 2018; Aicher et al. 2020], with a truncation at $N_{iter} \ll N_c$.

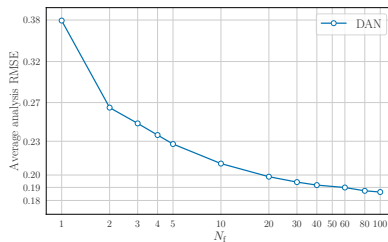
► For numerical efficiency, we choose to generate the samples *online*, as the training progresses, i.e. an *infinite training dataset*!

Hyperparameter sensitivity analysis

► Sensitivity analysis on key hyperparameters such as the number of trajectories N_T in the dataset, and the architecture parameters (N_f , N_b , N_{sb}) using the standard Lorenz 96 DA configuration ($\mathcal{H} = \mathbf{I}_x$, $\mathbf{R} = \mathbf{I}_x$).



Filtering performance
vs the number of trajectories N_T



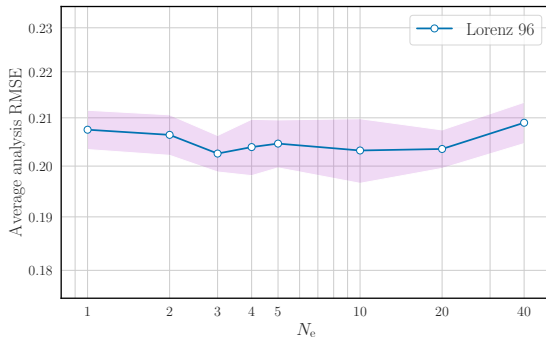
Filtering performance
vs the number of filters N_f

► The learned DA scheme *yields EnKF-like accuracy!*

► Compromise between α_θ 's size and its accuracy: $N_T = 2^{18}$, $N_f = 40$, $N_b = 5$, $N_{sb} = 5$.

Sensitivity to the ensemble size

► **First key observation:** The performance of a_θ barely depends on the ensemble size N_e . Hence localisation is irrelevant and unnecessary.



► **Second key observation:** a_θ does not require inflation and is incredibly robust to noise (as we shall see it applies its own inflation).

► **Explanation from the optimisation standpoint:** *feature collapse* of a_θ with respect to N_e in the training. Potential better solution when $N_e > 1$, but a_θ with $N_e = 1$ is as accurate as the EnKF!

Consequences and further checks

► Hence, from now on, we will focus on the mode: $N_e = 1$.

► Recall

$$\mathbf{E}_k^a = \mathbf{E}_k^f + a_\theta \left(\mathbf{E}_k^f, \mathbf{H}_k^\top \mathbf{R}_k^{-1} \delta_k \right). \quad (5)$$

► Performance of a_θ compared to baselines such as optimally tuned 3D-Var, the learned optimal linear filter, optimally tuned EnKF:

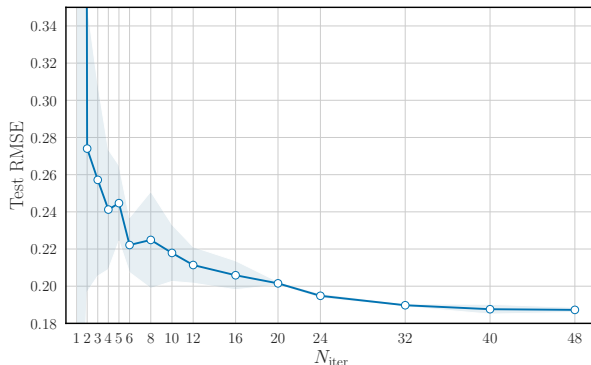
DA method	well-tuned classical	DL-based	aRMSE
EnKF-N, $N_e = 20$	yes		0.191
EnKF-N, $N_e = 40$	yes		0.179
3D-Var	yes		0.40
a_θ , $N_e = 1$, $N_f = 40$		yes	0.191
a_θ , $N_e = 1$, $N_f = 100$		yes	0.185
linear a_θ , $N_e = 1$, $N_f = 40$		yes	0.384
simplified $\hat{a}_\theta$, $N_e = 1$, $N_f = 40$		yes	0.382

where (simplified Ansatz)

$$\mathbf{E}_k^a = \mathbf{E}_k^f + \hat{a}_\theta \left(\mathbf{H}_k^\top \mathbf{R}_k^{-1} \delta_k \right). \quad (6)$$

Sensitivity to the training (truncation) depth N_{iter}

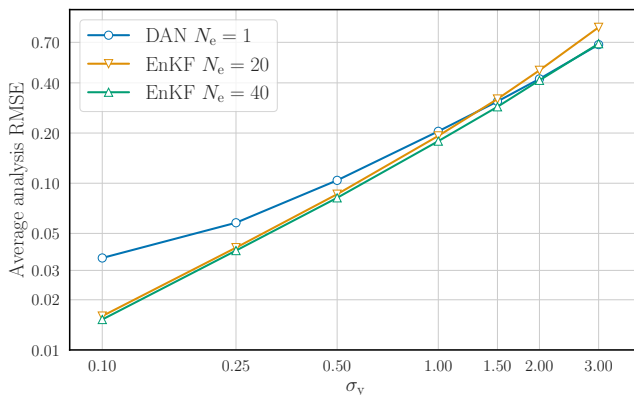
- ▶ Training through $N_{\text{iter}} = 1$ cycle cannot learn about the direct impact of the dynamics on DA.
- ▶ Training through N_{iter} chained cycles is expected to be crucial to the accuracy and robustness of the learned a_θ [Bocquet et al. 2025].



- ▶ Training depth does matter!
As expected, $N_{\text{iter}} \geq 2$ cycles are required to see significant benefits.

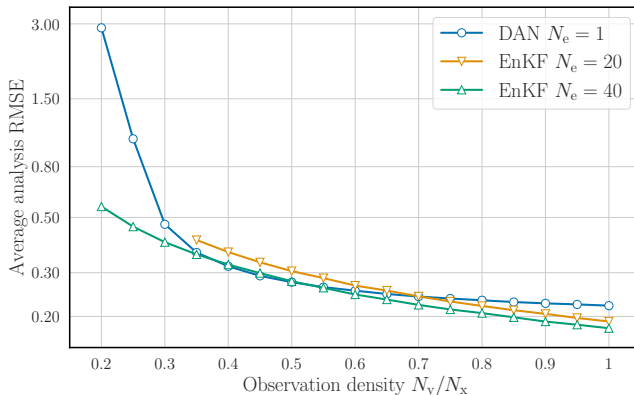
Sensitivity to observation error magnitude

- ▶ Next, we carry out a series of experiments that are not central to our message here but further ground the *viability of such learned a_θ* (assuming here $\mathbf{R}_k \stackrel{\Delta}{=} \mathbf{I}_x$).
- ▶ Impact of the *observation noise magnitude* on the data assimilation tests:



Sensitivity to observation sparsity

► Impact of the *sparsity of the observation dataset* on the data assimilation tests:



► a_θ trained with time-dependent, random, observation numbers and positions.

Operator expansion of the analysis

► We look for *a classical Kalman update that would be a good match to a_θ* seen as a mathematical map, at least for small analysis increments.

► To that end, we define the time-dependent normalised scalar anomalies

$$b_k = \frac{1}{\sqrt{N_x}} \|a_\theta(\mathbf{x}_k, \mathbf{0})\|, \quad (7)$$

along with the associated mean bias b and the standard deviation s of b_k in time.

→ We obtain $b \simeq 5 \times 10^{-3}$ and $s \simeq 10^{-3}$, which are indeed very small compared to the typical aRMSE of an either DAN or EnKF run, i.e., 0.20.

► Next, expanding with respect to the innovation, the following functional form for a_θ is assumed:

$$a_\theta(\mathbf{x}, \mathbf{H}^\top \mathbf{R}^{-1} \delta) \approx \mathbf{K}(\mathbf{x}) \cdot \delta, \quad (8)$$

owing to the fact that no state update is needed when the innovation vanishes, and only keeping the leading order term in δ .

Identifying the operators in the expansion

► Innovations $\{\delta_j\}_{j=1,\dots,N_p}$ are sampled from $\delta_j \sim N(\mathbf{0}, \mathbf{R})$.

This yields a set of corresponding incremental updates $\{\mathbf{a}_j = a_\theta(\mathbf{x}, \mathbf{H}^\top \mathbf{R}^{-1} \delta_j)\}_{j=1,\dots,N_p}$.

$\mathbf{K}(\mathbf{x})$ is then estimated with the least squares problem

$$\mathcal{L}_{\mathbf{x}}(\mathbf{K}) = \sum_{j=1}^{N_p} \left\| \mathbf{a}_j - \bar{\mathbf{a}} - \mathbf{K}(\mathbf{x}) \cdot (\delta_j - \bar{\delta}) \right\|^2, \quad (9)$$

where $\bar{\mathbf{a}} = N_p^{-1} \sum_{j=1}^{N_p} \mathbf{a}_j$ and $\bar{\delta} = N_p^{-1} \sum_{j=1}^{N_p} \delta_j$.

► Within the *best linear unbiased estimator* framework, $\mathbf{K}$ is related to $\mathbf{P}^a$ through $\mathbf{K} = \mathbf{P}^a \mathbf{H}^\top \mathbf{R}^{-1}$ so that from Eq. (8),

$$a_\theta(\mathbf{x}, \mathbf{H}^\top \mathbf{R}^{-1} \delta) \approx \mathbf{P}^a \mathbf{H}^\top \mathbf{R}^{-1} \delta, \quad (10)$$

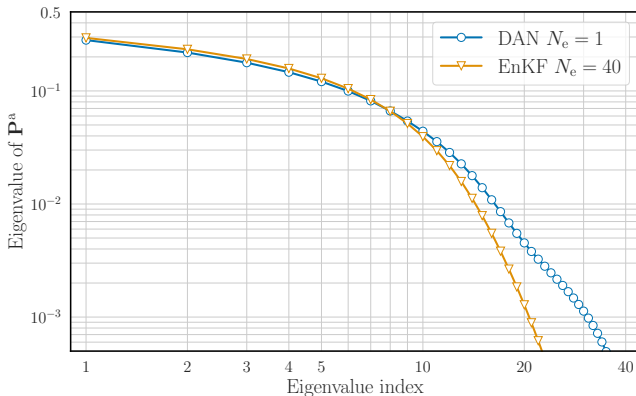
which suggests that an expansion in the second variable $\zeta \in \mathbb{R}^{N_x}$ of a_θ yields

$$a_\theta(\mathbf{x}, \zeta) \approx \mathbf{P}^a(\mathbf{x}) \cdot \zeta. \quad (11)$$

Hence, we can obtain a numerical estimation of an equivalent $\mathbf{P}^a(\mathbf{x})$.

What is learned? Supporting numerical results

- The surrogate $\mathbf{P}^a$, denoted $\mathbf{P}_{\text{DAN}}^a$ and estimated from Eq. (11), is compared to that of a concurrent well-tuned EnKF with $N_e = 40$, whose analysis error covariance matrix is $\mathbf{P}_{\text{EnKF}}^a$.



- Time-averaged eigenspectra of $\mathbf{P}_{\text{DAN}}^a$ and $\mathbf{P}_{\text{EnKF}}^a$. They are remarkably close to each other for the first 10 modes. Beyond these modes the a_θ operator is likely to selectively apply *some multiplicative inflation*, as one would expect from such stable DA runs.

Main interpretation

- **Conclusion 1:** a_θ depends on the innovation but also directly on $\mathbf{x}_k^f$ when $N_e = 1$, as opposed to the incremental update of the EnKF: *a_θ extracts important information from $\mathbf{x}_k^f$.*
- **Conclusion 2:** a_θ manages to assess a $\mathbf{P}_{\text{DAN}}^a$ with $N_e = 1$ which is very close to $\mathbf{P}_{\text{EnKF}}^a$ with $N_e = 40$, for the dominant axes, and applies multiplicative inflation on the less unstable modes.² We conclude that *a_θ directly learns about the dynamics features.* Hence, for a_θ , critical pieces of information on $\mathbf{P}_k^a$ are encoded, and thus exploitable, in $\mathbf{x}_k^f$ alone.
 —→ Supported by results from [Sacco et al. 2024](#); [Sakov 2025](#).
- **Explanation, conclusion 3:** Furthermore, if the DA run (the forecast and analysis cycle) is considered as an ergodic dynamical system of its own,³ the *multiplicative ergodic theorem* guarantees the existence of a mapping between $\mathbf{x}_k^f$ and $\mathbf{P}_k^a$ that a_θ is able to guess. We believe that a generalised variant of *the multiplicative ergodic theorem for non-autonomous random dynamics should be applicable*.⁴

²[Bocquet et al. 2015]

³[Carrassi et al. 2008]

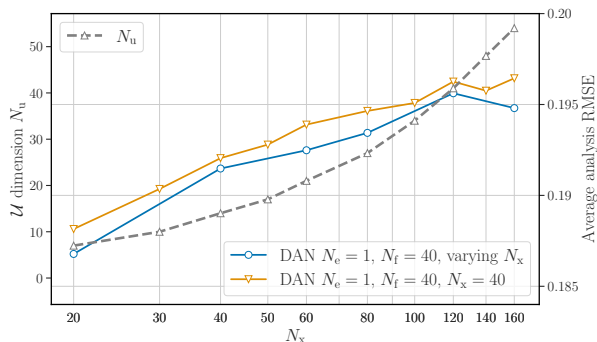
⁴[Arnold 1998; Flandoli et al. 2021]

Outline

- 1 Learning data assimilation from artificial intelligence: first results
 - Sequential data assimilation for chaotic dynamics
 - Learning data assimilation
 - Sensitivity analysis and first results
 - Investigation and interpretation
- 2 Learning data assimilation from artificial intelligence: advanced results
 - Scalability and locality
 - Learning the analysis operator from observations only
 - Learning nonlinearity
- 3 Conclusions

Locality and scalability – 1/2

► a_θ is now trained without changing the architecture and the hyperparameters ($N_f = 40$), but with a changing state space dimension $N_x \in [20, 160]$. Almost as good as well tuned EnKFs with changing dimension N_x and $N_e = N_x$!



→ We conjecture that a_θ extracts *local* pieces of information from $\mathbf{x}_k^f$.

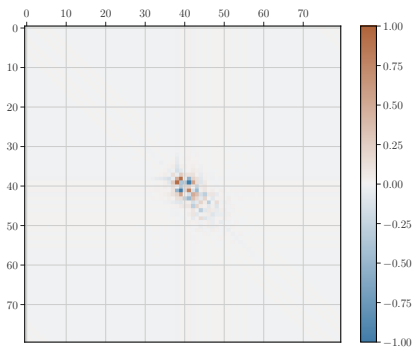
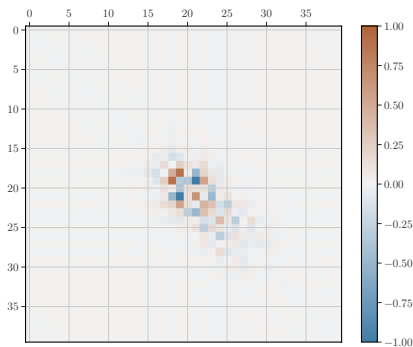
► a_θ , learned from Lorenz 96 with $N_x = 40$ is now tested on Lorenz 96 models with N_x ranging from 20 to 160 (same weights and biases!). The performance is still on par with retraining! We called this a *transdimensional transfer*.

Locality and scalability – 2/2

► These local patterns (for a_θ , not $\mathcal{M}$) can be pictured from the sensitivity:

$$\mathbf{S} = \left\langle \mathbf{C} : \left[\nabla_{\mathbf{x}} \nabla_{\zeta} a_\theta(\mathbf{x}, \zeta) |_{\zeta=0} \right] \right\rangle_{\mathbf{x} \in \mathcal{T}} = \left\langle \mathbf{C} : [\nabla_{\mathbf{x}} \mathbf{P}^a(\mathbf{x})] \right\rangle_{\mathbf{x} \in \mathcal{T}}, \quad (12)$$

where $\mathcal{T}$ is a long L96 trajectory, and $\mathbf{C}$ is a tensor that leverages translational invariance of the L96 model: $[\mathbf{C}]_{ij}^{nmk} = \frac{1}{N_{\mathbf{x}}} \delta_{n,i+k} \delta_{m,j+k}$.



Semi-supervised learning

- What if we do not have access to the truth $\mathbf{x}_k^t$ but to the observations only $\mathbf{y}_k$?
- Assume (i) $\mathcal{H}_k$ is linear, (ii) $\mathbf{y}_k \perp \mathbf{y}_{k+1}$, and the estimator $\mathbf{z}_{k+1}^\theta$ only depends on $(\mathbf{x}_k, \mathbf{y}_k)$.
- We define the **semi-supervised** loss function as

$$\mathcal{L}(\theta) = \sum_{k=1}^{N_c} \left\| \mathbf{y}_k - \mathbf{H}_k \mathbf{z}_k^\theta \right\|^2 = \sum_{k=1}^{N_c} \mathcal{L}_k(\theta). \quad (13)$$

But we have from the above assumptions:

$$\mathbb{E}_{\mathbf{y}} [\mathcal{L}_k(\theta)] = \mathbb{E}_{\mathbf{y}} \left[\left\| \mathbf{y}_k - \mathbf{H}_k \mathbf{x}_k^t \right\|^2 \right] + \mathbb{E}_{\mathbf{y}} \left[\left\| \mathbf{H}_k (\mathbf{x}_k^t - \mathbf{z}_k^\theta) \right\|^2 \right] \quad (14a)$$

$$= \text{Cst} + \mathbb{E}_{\mathbf{y}} \left[\left\| \mathbf{H}_k (\mathbf{x}_k^t - \mathbf{z}_k^\theta) \right\|^2 \right] \quad (14b)$$

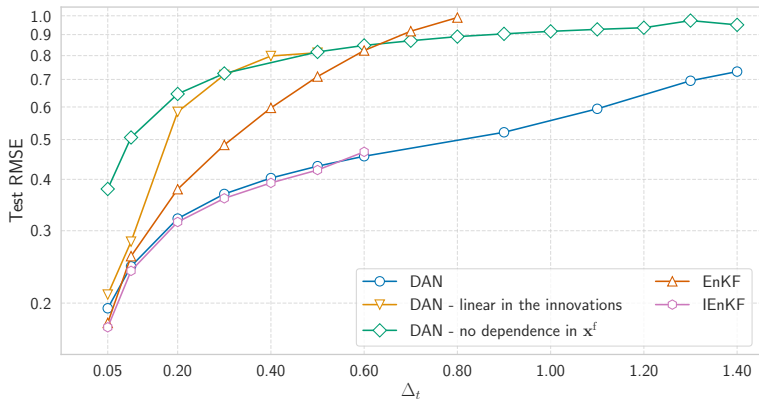
- Hence, generalising [McCabe et al. 2021] to non-trivial $\mathbf{H}_k$, we can learn $\mathbf{z}_k^\theta$ from the observation only, with further assumptions on $\{\mathbf{H}_k\}_{k=1, \dots, K}$. For instance, we can choose $\mathbf{z}_k^\theta$ such that:⁵

$$\mathcal{L}_k(\theta) = \left\| \mathbf{y}_{k+1} - \mathbf{H}_{k+1} \mathcal{M} \left\{ \mathbf{x}_k^f + a_\theta \left(\mathbf{x}_k^f, \mathbf{H}_k^\top \mathbf{R}_k^{-1} (\mathbf{y}_k - \mathbf{H}_k \mathbf{x}_k^f) \right) \right\} \right\|^2. \quad (15)$$

⁵[Bocquet et al. 2025]

Data assimilation networks in stronger nonlinear conditions

- ▶ Testing DANs as *the update time-step Δ_t is increased*, with the L96 model.
- ▶ Comparison with well tuned *EnKF*, *IEnKF*, and well-tuned static background DA methods.



- ▶ Performing at least as well as the IEnKF, *without an ensemble, without any nonlinear iterative solver, without inflation and localisation!*
- ▶ In addition to the MET map, DANs also implicitly *learn non-Gaussian priors*.

Outline

- 1 Learning data assimilation from artificial intelligence: first results
 - Sequential data assimilation for chaotic dynamics
 - Learning data assimilation
 - Sensitivity analysis and first results
 - Investigation and interpretation
- 2 Learning data assimilation from artificial intelligence: advanced results
 - Scalability and locality
 - Learning the analysis operator from observations only
 - Learning nonlinearity
- 3 Conclusions

Conclusions

- ▶ One can learn *robust* DA methods, *without an ensemble*, *without inflation*, that are *as accurate* as the very best baseline methods (EnKF), including in *stronger nonlinear regimes* (IEnKF).
- ▶ We have carried similar numerical experiments with the *Kuramoto–Sivashinski* model and a *single-layer QG model on the sphere*, with similar conclusions.
- ▶ The performance is achieved through (i) implicitly learning the map $\mathbf{x}^f \mapsto \mathbf{P}^f(\mathbf{x}^f)$, and (ii) through implicitly learning non-Gaussian priors!
- ▶ This suggests that end-to-end approaches such as GraphDOP that only have a snapshot of the physical system, can still implicitly rely on dynamical and non-Gaussian priors!
- ▶ Will such *multiplicative ergodic theorem* still be valid in more anisotropic, non-autonomous, forced, multivariate, heterogeneously observed systems?
- ▶ *In any case, this promotes a rethinking of the popular sequential DA schemes for chaotic dynamics.*

Talk mainly based on:

- ▶ M. Bocquet, A. Farchi, T. S. Finn, C. Durand, S. Cheng, Y. Chen, I. Pasmans, and A. Carrassi. "Accurate deep learning-based filtering for chaotic dynamics by identifying instabilities without an ensemble". In: Chaos 29 (2024), p. 091104.
- ▶ M. Bocquet, T. S. Finn, S. Cheng, W. Yu, and A. Farchi. "On the performance of data assimilation neural networks in nonlinear conditions". In: (2025). In preparation.

References |

- [1] C. Aicher, N. J. Foti, and E. B. Fox. "Adaptively Truncating Backpropagation Through Time to Control Gradient Bias". In: *Proceedings of The 35th Uncertainty in Artificial Intelligence Conference*. Ed. by Ryan P. Adams and Vibhav Gogate. Vol. 115. Proceedings of Machine Learning Research. PMLR, 22–25 Jul 2020, pp. 799–808.
- [2] M. Alexe et al. *GraphDOP: Towards skilful data-driven medium-range weather forecasts learnt and initialised directly from observations*. 2024. arXiv: 2412.15687 [physics.ao-ph].
- [3] A. Allen et al. "End-to-end data-driven weather prediction". In: *Nature* (2025).
- [4] L. Arnold. *Random Dynamical Systems*. Springer Berlin, Heidelberg, 1998, p. 586.
- [5] M. Asch, M. Bocquet, and M. Nodet. *Data Assimilation: Methods, Algorithms, and Applications*. Fundamentals of Algorithms. SIAM, Philadelphia, 2016, p. 324.
- [6] M. Bocquet, A. Farchi, T. S. Finn, C. Durand, S. Cheng, Y. Chen, I. Pasmans, and A. Carrassi. "Accurate deep learning-based filtering for chaotic dynamics by identifying instabilities without an ensemble". In: *Chaos* 29 (2024), p. 091104.
- [7] M. Bocquet, T. S. Finn, S. Cheng, W. Yu, and A. Farchi. "On the performance of data assimilation neural networks in nonlinear conditions". In: (2025). in preparation.
- [8] M. Bocquet, P. N. Raanes, and A. Hannart. "Expanding the validity of the ensemble Kalman filter without the intrinsic need for inflation". In: *Nonlin. Processes Geophys.* 22 (2015), pp. 645–662.
- [9] M. Bocquet and P. Sakov. "Joint state and parameter estimation with an iterative ensemble Kalman smoother". In: *Nonlin. Processes Geophys.* 20 (2013), pp. 803–818.
- [10] P. Boudier, A. Fillion, S. Gratton, S. Gürol, and S. Zhang. "Data Assimilation Networks". In: *J. Adv. Model. Earth Syst.* 15 (2023), e2022MS003353.
- [11] A. Carrassi, M. Ghil, A. Trevisan, and F. Uboldi. "Data assimilation as a nonlinear dynamical systems problem: Stability and convergence of the prediction-assimilation system". In: *Chaos* 18 (2008), p. 023112.
- [12] Y. Chen, D. Sanz-Alonso, and R. Willett. "Autodifferentiable Ensemble Kalman Filters". In: *SIAM J. Math. Data Sci.* 4 (2022), pp. 801–833.
- [13] R. S. Cintra and H. F. de Campos Velho. "Data assimilation by artificial neural networks for an atmospheric general circulation model". In: *Advanced applications for artificial neural networks*. Ed. by A. ElShahat. IntechOpen, 2018. Chap. 17, pp. 265–286.
- [14] F. Flandoli and E. Tonello. *An introduction to random dynamical systems for climate*. 2021.
- [15] T. Haarnoja, A. Ajay, S. Levine, and P. Abbeel. "Backprop KF: Learning Discriminative Deterministic State Estimators". In: *Advances in Neural Information Processing Systems*. Ed. by D. Lee, M. Sugiyama, U. Luxburg, I. Guyon, and R. Garnett. Vol. 29. Curran Associates, Inc., 2016.

References II

- [16] T. P. Härter and H. F. de Campos Velho. "Data Assimilation Procedure by Recurrent Neural Network". In: *Eng. Appl. Comput. Fluid Mech.* 6 (2012), pp. 224–233.
- [17] H.S. Hoang, P. De Mey, and O. Talagrand. "A simple adaptive algorithm of stochastic approximation type for system parameter and state estimation". In: *Proceedings of 1994 33rd IEEE Conference on Decision and Control*. Vol. 1. 1994, 747–752 vol.1.
- [18] S. Hoang, R. Baraille, O. Talagrand, X. Carton, and P. De Mey. "Adaptive filtering: application to satellite data assimilation in oceanography". In: *Dynam. Atmos. Ocean* 27 (1998), pp. 257–281.
- [19] P. Lean, M. Alexe, E. Boucher, E. Pinnington, S. Lang, P. Laloyaux, N. Bormann, and A. McNally. *Learning from nature: insights into GraphDOP's representations of the Earth System*. 2025. arXiv: 2508.18018 [physics.ao-ph].
- [20] E. Luk, E. Bach, R. Baptista, and A. Stuart. *Learning Optimal Filters Using Variational Inference*. 2024. arXiv: 2406.18066 [cs.LG].
- [21] M. McCabe and J. Brown. "Learning to Assimilate in Chaotic Dynamical Systems". In: *Advances in Neural Information Processing Systems*. Ed. by M. Ranzato, A. Beygelzimer, Y. Dauphin, P.S. Liang, and J. Wortman Vaughan. Vol. 34. Curran Associates, Inc., 2021, pp. 12237–12250.
- [22] M. A. Sacco, M. Pulido, J. J. Ruiz, and P. Tandeo. "On-line machine-learning forecast uncertainty estimation for sequential data assimilation". In: *Q. J. R. Meteorol. Soc.* 150 (2024), pp. 2937–2954.
- [23] P. Sakov. *On building the state error covariance from a state estimate*. 2025. arXiv: 2411.14809 [nlin.CD].
- [24] H. Tang and J. Glass. "On Training Recurrent Networks with Truncated Backpropagation Through time in Speech Recognition". In: *2018 IEEE Spoken Language Technology Workshop (SLT)*. 2018, pp. 48–55.