

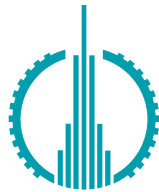
Nov 2025 - Apr 2026

Internship thesis

NONLINEAR COMPRESSED REDUCED BASIS METHOD FOR DATA ASSIMILATION



Université
de Rennes



巴黎卓越工程师学院
Ecole d'Ingénieurs Paris SJTU



上海交通大学

SHANGHAI JIAO TONG UNIVERSITY

Author
Raphaël LECOQ

Supervised by
Helin GONG, Yvon MADAY

Abstract

Reduced-order models are widely used to accelerate the simulation and assimilation of parametrized partial differential equations, but their performance is fundamentally limited by the linear approximation capabilities of classical reduced basis methods. This work proposes a nonlinear compressed reduced basis framework for data assimilation that combines Proper Orthogonal Decomposition (POD), the Empirical Interpolation Method (EIM), and machine learning. We first revisit the theoretical foundations of reduced basis methods, establishing the connection between POD approximation errors and the Kolmogorov N-width of the solution manifold. Motivated by the observation that the intrinsic dimension of the parameter manifold is often much smaller than the number of POD modes required for accurate reconstruction, we decompose the reduced coordinates into a low-dimensional linear component and a nonlinear residual. The latter is approximated by a Gaussian Process Regression model that learns the mapping between low order and high order POD coefficients. This nonlinear closure is incorporated into several data assimilation strategies, including direct optimization, gradient-based reconstruction, and fixed-point iterations, to recover the full state from sparse and noisy sensor observations. Numerical experiments on analytical benchmark problems and nonlinear transport equations demonstrate that the proposed approach accurately learns the latent nonlinear structure of the reduced coordinates and improves state reconstruction while reducing the number of optimization variables. The framework aims to provide a flexible and computationally efficient extension of classical POD-EIM techniques for nonlinear reduced-order modeling and sensor-based data assimilation.

Acknowledgment :

This internship was made possible thanks to Helin GONG¹. I am sincerely grateful to him for welcoming me to Shanghai and into his research team at the Shanghai Paris Elite Institute of Technology. His guidance, support, and availability throughout this internship were invaluable.

I would also like to express my sincere thanks to Yvon MADAY² for his supervision from Paris and for the many insightful discussions we had during the internship. His advice and perspective greatly contributed to this work.

Most of the ideas presented in this report originated from our joint discussions. I am especially grateful for the time both Helin GONG and Yvon MADAY devoted to understanding my results, helping me interpret them, and suggesting new directions. Their scientific insight and encouragement were instrumental in shaping this work.

Finally, I would like to warmly thank Xinyue GUO³. When I arrived in Shanghai, she had already laid the foundations of the project and established an efficient workflow. Throughout the internship, we worked closely together to develop the code and produce the results presented in this report. Her dedication, technical skills, and collaborative spirit made this project both productive and enjoyable.

Find the code on Github: https://github.com/RaphaelLcq/SJTU_NL_ROM_project.

For any questions, typo or corrections:

raphael.lecoq@ens-rennes.fr

¹Associate Professor, Paris Elite Institute of Technology, Shanghai Jiao Tong University, Shanghai, China.

²Professor, Sorbonne University, Laboratoire Jacques-Louis Lions, Paris, France.

³Bachelor student, Paris Elite Institute of Technology, Shanghai Jiao Tong University, Shanghai, China.

Contents

Part 1 Reduced Basis Method

I -	Approximation of the solution manifold	4
	1) Physic state	4
	2) Approximation by a reduced basis	5
II -	Generation of the Reduced Basis by Proper Orthogonal Decomposition	7
III -	Error modeling	10
	1) Model error	10
	2) Decomposition in linear and nonlinear basis	12

Part 2 Nonlinear compressed reduced basis method for EIM

I -	Sensor-observed system	13
	1) Observation function	13
	2) Field reconstruction from partial observations	14
	3) Succinct analysis of the noisy EIM error	16
	4) Data assimilation methods	17
	5) Learning the mapping between coefficients	19
II -	Numerical results	20
	1) Test cases	20
	2) Orthogonal modes results	21
	3) Advection of a gaussian bump	26
	Conclusion	33

Part 3 Appendix

A	Error metrics	34
	1) Parameter norm	34
	2) Timespace norm	34
	3) Error metrics	35
B	Test cases construction	36
	1) Orthogonal modes	36
	2) Accelerating advection of a gaussian bump	36
	3) Burgers equation	37
	List of tables	41
	List of figures	42
	Bibliography	45

Part 1

Reduced Basis Method

I - Approximation of the solution manifold

1) Physic state

Let X be a Hilbert⁴ of functions over a space-time domain $\Omega \times]0; T[$, with $\Omega \subset \mathbb{R}^d$ ($d \geq 1$) open⁵ space domain, endowed with the inner product $\langle \cdot | \cdot \rangle_X$ and its induced norm $\|\cdot\|_X$. Let $P \in \mathbb{N}^*$, $\mathbb{P} \subset \mathbb{R}^P$ be a compact parameter space. Let $F : X \times \mathbb{P} \mapsto \mathbb{R}$ a functional.

We consider the physical field parameterized by $\mu \in \mathbb{P}$:

$$\text{Find } u = u(\cdot, \cdot; \mu) \in X, \quad F(u, \mu) = 0. \quad (1.1)$$

In what follows we suppose the physic system admits an unique solution $u(\cdot, \cdot; \mu)$ for each $\mu \in \mathbb{P}$. Such physical system usually derives from a parameterized partial differential equation (PPDE):

Example 1: Parameterized partial differential equation [RAU91; HRS16]

Let $L_\mu : X \rightarrow X'$ be a linear operator that includes the linear differential operators^a.

Let $\mathcal{G}_\mu : X \mapsto X'$ be a nonlinear operator that includes the nonlinear differential operators^b.

A *Parameterized Partial Differential Equation* can be written :

$$\text{Find } u = u(\cdot, \cdot; \mu) \in X, \quad \begin{cases} \frac{\partial u}{\partial t} + L_\mu(u) + \mathcal{G}_\mu(u) = 0, \\ \text{Initial conditions,} \\ \text{Boundary conditions,} \\ \text{Physical constraints.} \end{cases} \quad (1.2)$$

Here the PPDE is supposed to have (in the time-dependant case) only order one time derivative, a linear part L_μ and a nonlinear part $\mathcal{G}_\mu$.

^aSee [LAN96; KYT96] for definition.

^bSee [LE 13; TAY91; KYT96] for definition.

We define the manifold described in X by the solution of (1.1) when the parameter μ describes $\mathbb{P}$:

Definition 1.1: Solution manifold

We define the *solution manifold* [RHE88] as:

$$\mathcal{M} := \{u(\cdot, \cdot; \mu) / \mu \in \mathbb{P}, \quad F(u, \mu) = 0\} \subset X.$$

This is the quantity of interest that we want to approximate in what follow.

Such system of equations is hard to describe in a closed analytical form, thus the use of numerical methods to obtain an approximation of the solution. Some methods include Galerkin elements methods [EG04; DE12; FIL24; QSS06].

In practice, these methods discretize the space X (in space and time) into a linear space $\mathbf{X}_\delta$ of finite dimension denoted $n_d \geq 1$.

⁴See [BRE11] for functional analysis theory.

⁵See [MEN90] for topology theory.

In the time-dependent case, $\mathbf{u}_{n_d} \in \mathbb{R}^{n_d}$ will be isomorph to a matrix where the coefficient u_{ij} corresponds to the i -th spatial degree of freedom at the j -th timestep.

For the sake of simplicity we assume that we are working in a conforming Galerkin framework [AG21, Section 18 and 19], where $\mathbf{X}_\delta \subset X$ such that $F(\mathbf{u}_{n_d}, \boldsymbol{\mu})$ is well defined. For $\boldsymbol{\mu} \in \mathbb{P}$, the *truth problem* then reads :

$$\text{Find } \mathbf{u}_{n_d} = \mathbf{u}_{n_d}(\boldsymbol{\mu}) \in \mathbf{X}_\delta, \quad F(\mathbf{u}_{n_d}, \boldsymbol{\mu}) = 0. \quad (1.3)$$

We call *truth solution* a solution $\mathbf{u}_{n_d}$ of the truth problem. From the truth problem, one can construct the truth solution manifold :

Definition 1.2: Truth solution manifold

We define the *truth solution manifold* as

$$\mathcal{M}_\delta := \{\mathbf{u}_{n_d}(\boldsymbol{\mu}) / \boldsymbol{\mu} \in \mathbb{P}, F(\mathbf{u}_{n_d}, \boldsymbol{\mu}) = 0\} \subset \mathbf{X}_\delta.$$

For $\delta > 0$ and for n_d big enough^a it is expected that :

$$\forall \boldsymbol{\mu} \in \mathbb{P}, \|\mathbf{u}_{n_d}(\boldsymbol{\mu}) - u(\boldsymbol{\mu})\|_X \leq \delta. \quad (1.4)$$

^aStrongly depends on the numerical method, e.g. for order 1 Lagrange finite elements $n_d \propto \frac{1}{\delta}$, [EG04; QSS06].

Remark :

The issue with the truth problems is the computational complexity, e.g. for finite elements it usually scales like $(n_d)^d$ [HRS16, Section 3.5].^a In particular, this not tractable for high precision simulations or high dimensional PDEs.

^aRecall that $d \leq 1$ is the spatial dimension and n_d is the discrete space-time dimension.

In the case of conforming Galerkin finite element methods, Cea's lemma [CEA64] give an a priori error control :

$$\exists C > 0, \forall \boldsymbol{\mu} \in \mathbb{P}, \quad \|u(\boldsymbol{\mu}) - \mathbf{u}_{n_d}(\boldsymbol{\mu})\|_X \leq C \inf_{\mathbf{v}_{n_d} \in \mathbf{X}_\delta} \|u(\boldsymbol{\mu}) - \mathbf{v}_{n_d}\|_X. \quad (1.5)$$

The approximation is said to be *quasi-optimal* if it satisfies an estimation equivalent to Eq.(1.5). In what follow, we suppose we have such property of quasi-optimality of the truth problem.

2) Approximation by a reduced basis

To answer the computational issue we consider a method with reduced linear dimension. We look for a reference function $\mathbf{u}_{\text{ref}} \in \mathbf{X}_\delta$ and a reduced basis $\mathcal{B}_N = [\boldsymbol{\zeta}^1, \dots, \boldsymbol{\zeta}^N] \in (\mathcal{M}_\delta - \mathbf{u}_{\text{ref}})^N$, $N \geq 1$ such that

$$\mathbf{V}_N := \text{Span} \{(\boldsymbol{\zeta}^i)_{1 \leq i \leq N}\} + \mathbf{u}_{\text{ref}} \subset \mathcal{M}_\delta$$

is a reduced space of dimension N that aims to represent $\mathcal{M}_\delta$ with high fidelity and dimension $N \ll n_d$.

We call the linear space $\mathbf{V}_N$ a *Reduced Basis* (RB) of the PPDE. Its goal is to approximate $\mathcal{M}_\delta$ as well as possible while having the smallest dimension allowing it.

Such method of constructing a RB to approximate a infinite-dimensional solution manifold is the so called *Reduced Basis Method* [OS16].

Searching an approximation of the physical field $u \in X$ by $\mathbf{u}_N \in \mathbf{V}_N$ then read as :

$$\text{For } \boldsymbol{\mu} \in \mathbb{P}, \text{ find } \mathbf{u}_N(\boldsymbol{\mu}) = \mathbf{u}_{\text{ref}} + \sum_{i=1}^N \hat{u}_i(\boldsymbol{\mu}) \boldsymbol{\zeta}^i \in \mathbf{V}_N, \quad F(\mathbf{u}_N, \boldsymbol{\mu}) = 0. \quad (1.6)$$

The formulation of the equation is the same as before yet the computational cost of the reduced basis is of the order $\mathcal{O}(N^d)$ which can be several amplitudes lesser than $\mathcal{O}(n_d)^d$.

Since the reduced basis is a linear subset of $\mathbf{X}_\delta$, we recover quasi-optimality [HRS16, Eq.(3.7)] :

$$\exists C > 0, \forall \boldsymbol{\mu} \in \mathbb{P}, \quad \|u(\boldsymbol{\mu}) - \mathbf{u}_N(\boldsymbol{\mu})\|_X \leq C \inf_{\mathbf{v}_N \in \mathbf{V}_N} \|u(\boldsymbol{\mu}) - \mathbf{v}_N\|_X. \quad (1.7)$$

To compare $\mathbf{V}_N$ with $\mathcal{M}$ we consider their Hausdorff distance^{6,7} in the X -norm :

$$E(\mathcal{M}, \mathbf{V}_N ; \|\cdot\|_X) := \sup_{u \in \mathcal{M}} \inf_{\mathbf{v}_N \in \mathbf{V}_N} \|u - \mathbf{v}_N\|_X = \sup_{\boldsymbol{\mu} \in \mathbb{P}} \min_{\mathbf{v}_N \in \mathbf{V}_N} \|u(\boldsymbol{\mu}) - \mathbf{v}_N\|_X. \quad (1.8)$$

The $\|\cdot\|_X$ can be omitted if there's no ambiguity on the considered norm. The inf is attained because X is a Hilbert [BRE11, Theorem 5.2] and $\mathbf{V}_N$ a finite dimensional linear subset.

We then focus on the optimal distance one can expect with such a approximation of $\mathcal{M}$.

Definition 1.3: Kolmogorov N -width [KOL36]

The *Kolmogorov N -width* measures the optimal Hausdorff distance between the approximation of a manifold $\mathcal{M}$ and a N -dimensional linear space included in $\mathcal{M}$:

$$d_N(\mathcal{M} ; \|\cdot\|_X) := \inf_{\substack{\mathbf{V}_N \subset \mathcal{M} \\ \dim(\mathbf{V}_N) = N}} E(\mathcal{M}, \mathbf{V}_N ; \|\cdot\|_X).$$

The $\|\cdot\|_X$ can be omitted if there's no ambiguity on the considered norm.

Because of the orthogonality property Eq.(1.8) and denoting by $\Pi_{\mathbf{V}_N}$ the orthogonal projection from X into $\mathbf{V}_N$ we can rewrite this quantity as :

$$d_N(\mathcal{M}) = \inf_{\substack{\mathbf{V}_N \subset \mathcal{M} \\ \dim(\mathbf{V}_N) = N}} \|u - \Pi_{\mathbf{V}_N}(u)\|_{L^\infty(\mathbb{P}, \mathbf{V}_N)}. \quad (1.9)$$

Assuming the numerical solving method of the physical field is as precise as we want, we can write :

$$\forall \boldsymbol{\mu} \in \mathbb{P}, \|u(\boldsymbol{\mu}) - \mathbf{u}_{n_d}(\boldsymbol{\mu})\|_X \leq \delta \implies |d_N(\mathcal{M}_\delta) - d_N(\mathcal{M})| \leq \delta. \quad (1.10)$$

It is equivalent to estimate Kolmogorov N -width with either the truth solutions or the continuous ones. Thanks to Eq.(1.7) we get the following a priori control :

$$\exists C > 0, \forall \boldsymbol{\mu} \in \mathbb{P}, \|u(\boldsymbol{\mu}) - \mathbf{u}_N(\boldsymbol{\mu})\|_X \leq C d_N(\mathcal{M}). \quad (1.11)$$

In that sense, equivalence holds between the error of the reduced basis representation of the analytical solution and the optimal distance between the solution manifold and a reduced basis model.

PDE	Equation	Kolmogorov N -width decay	Reference
Diffusion/Heat equation	$\partial_t u + \partial_x(\mu \partial_x u) = f$	$d_N(\mathcal{M}) \lesssim \exp(-\alpha N)$	[MPT02]
Linear advection	$\partial_t u + \mu \partial_x u = f$	$d_N(\mathcal{M}) \gtrsim N^{-1/2}$	[OS16]
Wave equation	$\partial_t^2 u - \mu^2 \partial_x^2 u = f$	$d_N(\mathcal{M}) \gtrsim N^{-1/2}$	[GU19]
Stokes equation	$-\nu \Delta \mathbf{u} + (\mathbf{u} \cdot \nabla) \mathbf{u} + \nabla p = f$ $\operatorname{div}(\mathbf{u}) = 0$	$d_N(\mathcal{M}) \lesssim \exp(-N/3)$	[MS20]

Table 1: Kolmogorov N -width of some partial differential equations.

These estimations are only true for some specific situations since there the Kolmogorov N -width decay depends on the data of the PDE and the dimension of the parameter space [AGU23].

⁶[SEN90], note this is the full Hausdorff distance since $\mathbf{V}_N \subset \mathcal{M}_\delta$.

⁷Also called the *angle*, *deviation* or *discrepancy* between $\mathcal{M}$ and $\mathbf{V}_N$ [QMF16, Section 5.4].

II - Generation of the Reduced Basis by Proper Orthogonal Decomposition

A N -dimensional Proper Orthogonal Decomposition (POD) [LIA+02] space of X is a N -dimensional linear subset $\mathbf{V}_N \subset X$ that minimizes their distance in the least-square sense [LIA+02; BUL12; HRS16].

Let $N_s \geq 1$. Let $\mathbb{P}_h = \{\boldsymbol{\mu}^1, \dots, \boldsymbol{\mu}^{N_s}\} \subset \mathbb{P}$ be a discrete and finite point-set. Define a set of snapshots of the truth solutions of cardinality N approximating the truth manifold as :

$$\mathcal{M}_\delta(\mathbb{P}_h) := \{\mathbf{u}_{n_d}(\boldsymbol{\mu}) / \boldsymbol{\mu} \in \mathbb{P}_h\}.$$

Define the centered snapshots matrix $\mathbf{S} = [\mathbf{u}_{n_d}(\boldsymbol{\mu}^1) \dots \mathbf{u}_{n_d}(\boldsymbol{\mu}^{N_s})] \in \mathbb{R}^{n_d \times N_s}$. To construct the reduced basis $\mathbf{V}_N$ we apply a Singular Vector Decomposition (SVD) [LIA+02; BUL12] to $\mathbf{S}$.

Let $r = \text{rank}(\mathbf{S}) \leq \min(n_d, N)$. There exists $\mathbf{V} \in \mathbb{R}^{n_d \times r}$, $\mathbf{D} = \text{diag}(\sigma_1, \dots, \sigma_r) \in \mathbb{R}^{r \times r}$ and $\mathbf{W} \in \mathbb{R}^{N_s \times r}$ such that [BHK20] :

$$\mathbf{S} = \mathbf{V} \mathbf{D} \mathbf{W}^T, \quad \mathbf{V}^T \mathbf{V} = \mathbf{I}_r, \quad \mathbf{W}^T \mathbf{W} = \mathbf{I}_r. \quad (1.12)$$

The diagonal values of $\mathbf{D}$, $\sigma_i > 0$, are such that σ_i^2 is an eigenvalue of $\mathbf{S}^T \mathbf{S}$ and :

$$\sigma_1 \geq \dots \geq \sigma_r > 0. \quad (1.13)$$

Remark :

Some SVD [LIA+02; BUL12; QMF16] considers $\mathbf{U} \in \mathbb{R}^{n_d \times n_d}$, $\mathbf{D} \in \mathbb{R}^{n_d \times N_s}$ and $\mathbf{W} \in \mathbb{R}^{N_s \times N_s}$ and $p = \min(N_s, n_d)$ such that :

$$\sigma_1 \geq \dots \geq \sigma_r > 0 = \sigma_{r+1} = \dots = \sigma_p.$$

This holds no difference since the $p - r$ last columns of $\mathbf{V}$ are associated to $\ker(\mathbf{S}^T \mathbf{S})$.

The best N -rank approximation of $\mathbf{S}$ in the least square sense is [BHK20, Section 3.5] :

$$\mathbf{S}_N = \mathbf{V}_N \mathbf{D}_N \mathbf{W}_N^T. \quad (1.14)$$

where $\mathbf{V}_N \in \mathbb{R}^{n_d \times N}$ is obtained by the first N columns of $\mathbf{Q}$, $\mathbf{D}_N = \text{diag}(\sigma_1, \dots, \sigma_n) \in \mathbb{R}^{N \times N}$ and $\mathbf{W}_N \in \mathbb{R}^{N \times N}$ is obtained by the first N columns of $\mathbf{W}$.

Theorem 1.4: SVD error estimation

We have the following operator least square error [BHK20, Lemma 3.8 and Theorem 3.9] :

$$\|\mathbf{S} - \mathbf{S}_N\|_2^2 = \min_{\substack{\mathbf{A} \in \mathbb{R}^{n_d \times N} \\ \text{rank}(\mathbf{A})=N}} \|\mathbf{S} - \mathbf{A}\|_2^2 = \sigma_{N+1}^2. \quad \textcolor{blue}{a}$$

Schmidt-Eckart-Young theorem [QMF16, Theorem 6.1] gives the Frobenius least-square error :

$$\|\mathbf{S} - \mathbf{S}_N\|_F^2 = \min_{\substack{\mathbf{A} \in \mathbb{R}^{n_d \times N} \\ \text{rank}(\mathbf{A})=N}} \|\mathbf{S} - \mathbf{A}\|_F^2 = \sum_{k=N+1}^r \sigma_k^2. \quad \textcolor{blue}{b}$$

^aSpectral norm on linear operators of its argument [HJ12, Example 5.6.6].

^bFrobenius norm on matrices of its argument [HJ12, Eq.(5.6.0.2)].

Remark :

The notation with $r = \text{rank}(\mathbf{S})$ can be replaced by $p = \min(n_D, N_s)$ since p is the maximum rank that $\mathbf{S} \in \mathbb{R}^{n_d \times N_s}$ can admit and defining $\sigma_i = 0, \forall p \geq i > r$.
 Hence both notations for the SVD are equivalents, and the dimension of the SVD **must be** lower than r or $\mathbf{S}_N = \mathbf{S}$.

The SVD orthogonal basis is constructed by taking the N first left singular vectors of $\mathbf{S}$ i.e. the N first columns of $\mathbf{Q}$, that is $\mathbf{V}_N$.

We recall that the projection operator from $\mathbb{R}^{n_d}$ to $\text{Span}\{\mathbf{V}_N\}$ is :

$$\Pi_{\mathbf{V}_N} : \mathbf{x} \in \mathbb{R}^{n_d} \mapsto \mathbf{V}_N \mathbf{V}_N^T \mathbf{x}. \quad (1.15)$$

We have the following optimal property [QMF16, Proposition 6.1] :

$$\sum_{i=1}^N \|\mathbf{u}_{n_d}(\boldsymbol{\mu}^i) - \mathbf{V}_N \mathbf{V}_N^T \mathbf{u}_{n_d}(\boldsymbol{\mu}^i)\|_2^2 = \min_{\substack{\mathbf{A} \in \mathbb{R}^{n_d \times N} \\ \mathbf{A}^T \mathbf{A} = \mathbf{I}_N}} \sum_{i=1}^N \|\mathbf{u}_{n_d}(\boldsymbol{\mu}^i) - \mathbf{A} \mathbf{A}^T \mathbf{u}_{n_d}(\boldsymbol{\mu}^i)\|_2^2 \quad (1.16)$$

$$= \sum_{i=N+1}^r \sigma_i^2. \quad (1.17)$$

where here $\|\cdot\|_2$ refers to the euclidean 2-norm on $\mathbb{R}^{n_d}$ [HJ12, Eq.(5.2.1)].

Thus $\mathbf{V}_N$ columns are the N -optimal orthogonal vectors in the sense of the l^2 -projection [COL12, Section 6.2] from any snapshot taken in $\mathbf{S}$ to $\text{Span}(\mathbf{V}_N)$.

To get an optimal distance over the X -norm with elements of $\mathbf{X}_\delta$, we consider a symmetric definite positive matrix $\mathbf{M}_\delta$ such that for $\mathbf{u}_\delta, \mathbf{v}_\delta \in \mathbf{X}_\delta$, it holds $\langle \mathbf{u}_\delta | \mathbf{v}_\delta \rangle_X = \mathbf{u}_\delta^T \mathbf{M}_\delta \mathbf{v}_\delta$ and $\|\mathbf{v}_\delta\|_X^2 = \mathbf{v}_\delta^T \mathbf{M}_\delta \mathbf{v}_\delta$. In such framework, the orthogonal projections into a space spanned by a matrix $\mathbf{V}$ becomes :

$$\Pi_{\mathbf{V}_N}(\mathbf{x}) = \mathbf{V}_N \mathbf{V}_N^T \mathbf{M}_\delta \mathbf{x}. \quad (1.18)$$

Then by [QMF16, Proposition 6.2] and [HRS16, Section 3.2.1] :

Theorem 1.5: POD basis in the X -norm

Consider $\tilde{\mathbf{S}} = \mathbf{M}_\delta^{1/2} \mathbf{S}$, $\tilde{\mathbf{S}} = \tilde{\mathbf{V}} \tilde{\mathbf{D}} \tilde{\mathbf{W}}$ its full rank SVD and $\tilde{\mathbf{S}}_N = \tilde{\mathbf{V}}_N \tilde{\mathbf{D}}_N \tilde{\mathbf{W}}_N$ its N -dimensional reduced SVD.

Let $\mathbf{V}_N = \mathbf{M}_\delta^{1/2} \tilde{\mathbf{V}}_N$.

Then all the precedent results hold in the norm $\|\cdot\|_X$ (or $\|\cdot\|_{\mathbf{X}_\delta}$), in particular :

$$\begin{aligned} \sum_{i=1}^N \|\mathbf{u}_{n_d}(\boldsymbol{\mu}^i) - \mathbf{V}_N \mathbf{V}_N^T \mathbf{M}_\delta \mathbf{u}_{n_d}(\boldsymbol{\mu}^i)\|_X^2 &= \min_{\substack{\mathbf{A} \in \mathbb{R}^{n_d \times N} \\ \mathbf{A}^T \mathbf{M}_\delta \mathbf{A} = \mathbf{I}_N}} \sum_{i=1}^N \|\mathbf{u}_{n_d}(\boldsymbol{\mu}^i) - \mathbf{A} \mathbf{A}^T \mathbf{M}_\delta \mathbf{u}_{n_d}(\boldsymbol{\mu}^i)\|_X^2 \\ &= \sum_{i=N+1}^r \sigma_i^2. \end{aligned}$$

Since we are interested in a POD that approximates entirely $\mathcal{M}_\delta$ and not only a discrete subset $\mathcal{M}_\delta(\mathbb{P}_h)$, we consider the minimization problem [QMF16, Remark 6.3] and [HRS16, Eq.(3.6)] :

$$\delta_{N,2}(\mathcal{M}_\delta) = \min_{\substack{\mathbf{A} \in \mathbb{R}^{n_d \times N} \\ \mathbf{A}^T \mathbf{M}_\delta \mathbf{A} = \mathbf{I}_N}} \left(\int_{\mathbb{P}} \|\mathbf{u}_{n_d}(\boldsymbol{\mu}) - \mathbf{A} \mathbf{A}^T \mathbf{M}_\delta \mathbf{u}_{n_d}(\boldsymbol{\mu})\|_X^2 d\boldsymbol{\mu} \right)^{1/2}. \quad (1.19)$$

The minimum is attained [QMF16, Proposition 6.3] and Eq.1.19 is analogous to the Kolmogorov N -width for the $L^2(\mathbb{P}, \mathbf{X}_\delta)$ norm induced with $\|\cdot\|_X$ [QMF16, Remark 6.3] :

$$\delta_{N,2}(\mathcal{M}_\delta) = \min_{\substack{\mathbf{A} \in \mathbb{R}^{n_d \times N} \\ \mathbf{A}^T \mathbf{M}_\delta \mathbf{A} = \mathbf{I}_N}} \|\mathbf{u}_{n_d} - \mathbf{A} \mathbf{A}^T \mathbf{M}_\delta \mathbf{u}_{n_d}\|_{(L^2(\mathbb{P}, \mathbf{X}_\delta); \|\cdot\|_X)} \leq |\mathbb{P}|^{1/2} d_N(\mathcal{M}_\delta).$$

We need to bridge between the parameter-continuous framework and the parameter-discrete one. This is done [QMF16, Section 6.5].

Using quadrature formula we can write [QSS06, Section 9.1] :

$$\int_{\mathbb{P}} \|\mathbf{u}_{n_d}(\boldsymbol{\mu}) - \mathbf{V}_N \mathbf{V}_N^T \mathbf{M}_\delta \mathbf{u}_{n_d}(\boldsymbol{\mu})\|_X^2 d\boldsymbol{\mu} \simeq \sum_{i=1}^N w^i \|\mathbf{u}_{n_d}(\boldsymbol{\mu}^i) - \mathbf{V}_N \mathbf{V}_N^T \mathbf{M}_\delta \mathbf{u}_{n_d}(\boldsymbol{\mu}^i)\|_X^2. \quad (1.20)$$

where $w^i > 0$ are suitable quadrature weights and $\boldsymbol{\mu}^i \in \mathbb{P}$ are quadrature points in $\mathbb{P}$.

This approximation holds an error that we call *sampling error* :

$$\varepsilon_{\text{sampling error}}^2 := \left| \sum_{i=1}^N \left\| \mathbf{Q}^{1/2} \mathbf{u}_{n_d}(\boldsymbol{\mu}^i) - \mathbf{Q}^{1/2} \mathbf{V}_N \mathbf{V}_N^T \mathbf{M}_\delta \mathbf{u}_{n_d}(\boldsymbol{\mu}^i) \right\|_X^2 - \int_{\mathbb{P}} \|\mathbf{u}_{n_d}(\boldsymbol{\mu}) - \mathbf{V}_N \mathbf{V}_N^T \mathbf{M}_\delta \mathbf{u}_{n_d}(\boldsymbol{\mu})\|_X^2 d\boldsymbol{\mu} \right|.$$

The choice of the quadrature points depends on the dimension of the parameters space and define the sampling $\mathbb{P}_h = (\boldsymbol{\mu}^1, \dots, \boldsymbol{\mu}^{N_s})$ to construct $\mathcal{M}_\delta(\mathbb{P}_h)$.

We denote by $\mathbf{Q} = \text{diag}(w_1, \dots, w_N) \in \mathbb{R}^{N_s \times N_s}$ the diagonal weight matrix and define an SVD on $\tilde{\mathbf{S}} = \mathbf{M}_\delta^{1/2} \mathbf{S} \mathbf{Q}^{1/2}$, then the N first singular vectors $\mathbf{V}_N$ are such that [QMF16, Algorithm 6.3] :

$$\begin{aligned} & \sum_{i=1}^N w_i \|\mathbf{u}_{n_d}(\boldsymbol{\mu}^i) - \mathbf{V}_N \mathbf{V}_N^T \mathbf{M}_\delta \mathbf{u}_{n_d}(\boldsymbol{\mu}^i)\|_X^2 \\ &= \min_{\substack{\mathbf{A} \in \mathbb{R}^{n_d \times N} \\ \mathbf{A} \mathbf{M}_\delta \mathbf{A}^T = \mathbf{I}_N}} \left\| \tilde{\mathbf{S}} - \mathbf{A} \mathbf{A}^T \tilde{\mathbf{S}} \right\|_F^2 \\ &= \sum_{i=N+1}^r \sigma_i^2. \end{aligned}$$

Note the $\sigma_i > 0$ are the eigenvalues for the considered sampling. The approximation of $\mathcal{M}_\delta$ by a weighted POD is [QMF16, Eq.(6.35)] :

$$Cd_n(\mathcal{M}_\delta) \geq \delta_{N,2}(\mathcal{M}_\delta) \sim \varepsilon_{\text{sampling error}}^2 + \underbrace{\sum_{i=N+1}^r \sigma_i^2}_{\text{POD error}} + \varepsilon_{\text{numerical error}}^2 \geq d_N(\mathcal{M}_\delta). \quad (1.21)$$

For a parameter $\boldsymbol{\mu} \in \mathbb{P}$, we expect the error between the reduced solution $\mathbf{u}_N(\boldsymbol{\mu})$ and the truth solution $\mathbf{u}_{n_d}(\boldsymbol{\mu})$ to be equivalent to $d_N(\mathcal{M}_\delta)$ up to the sampling error $\varepsilon_{\text{sampling error}}$ and the numerical error $\varepsilon_{\text{numerical error}} \sim 10^{-12}$ [QSS06, Section 2.4].

Remark :

Let $N \geq 1$. Let $\mathbf{M} \in \mathbb{R}^{N \times N}$ be symmetric positive definite (SPD) matrix. Its norm induced by $\langle \mathbf{M} \mathbf{u} | \mathbf{u} \rangle = \|\mathbf{u}\|_{\mathbf{M}}^2$, $\mathbf{u} \in \mathbb{R}^N$, is such that if $\lambda_1 = \max \text{Sp}(\mathbf{M})$ and $\lambda_n = \min \text{Sp}(\mathbf{M})$:

$$\lambda_n \|\mathbf{u}\|_2^2 \leq \|\mathbf{u}\|_{\mathbf{M}}^2 \leq \lambda_1 \|\mathbf{u}\|_2^2.$$

Thus a POD on the 2-norm is equivalent in the sense of the norms to a POD on the X -norm if the matrix $\mathbf{M}_\delta$ is SPD.

Hence it is enough to do a sampled POD with the 2-norm in that case.

We conclude by the following statement that summarizes the method :

Theorem 1.6: Weighted POD in energy norm

Let $N_s \geq 1$, $(w_1, \dots, w_N) \in (\mathbb{R}_+^*)^{N_s}$ be quadrature weights and $(\boldsymbol{\mu}^1, \dots, \boldsymbol{\mu}^N) \in \mathbb{P}^{N_s}$ be quadrature points. Define $\mathbf{Q} = \text{diag}(w_1, \dots, w_N) \in \mathbb{R}^{N_s \times N_s}$.

Consider the matrix of truth solutions snapshots $\mathbf{S} = [\mathbf{u}_N(\boldsymbol{\mu}^1) \ \dots \ \mathbf{u}_N(\boldsymbol{\mu}^{N_s})] \in \mathbb{R}^{n_d \times N}$.

Let $\mathbf{M}_\delta \in \mathbb{R}^{n_d \times n_d}$ be the SPD matrix that defines $\mathbf{X}_\delta$ inner-product.

Let $\tilde{\mathbf{S}} = \mathbf{M}_\delta^{1/2} \mathbf{S} \mathbf{Q}^{1/2}$.

Then the rank N matrix $\mathbf{V}_N$ defined by the N first left singular vectors of $\tilde{\mathbf{S}}^T \tilde{\mathbf{S}}$ form a POD basis such that :

$$\begin{aligned} & \sum_{i=1}^N w_i \left\| \mathbf{u}_{n_d}(\boldsymbol{\mu}^i) - \mathbf{V}_N \mathbf{V}_N^T \mathbf{M}_\delta \mathbf{u}_{n_d}(\boldsymbol{\mu}^i) \right\|_X^2 \\ &= \min_{\substack{\mathbf{A} \in \mathbb{R}^{n_d \times N} \\ \mathbf{A} \mathbf{M}_\delta \mathbf{A}^T = \mathbf{I}_N}} \left\| \tilde{\mathbf{S}} - \mathbf{A} \mathbf{A}^T \tilde{\mathbf{S}} \right\|_F^2 = \sum_{i=N+1}^r \sigma_i^2. \end{aligned}$$

By Eq.(1.21) we note that the decrease rate of this method is directly linked to the Kolmogorov N -width, and if the latter has a very slow decay rate then a large number of mode is needed which also results in a large number of snapshots needed. However large number of modes can make POD-based ROMs slower than full order models.

III - Error modeling

1) Model error

Let $\mathbf{V}_N$ be a reduced space generated by POD. Recalling the form of the reduced solution $\mathbf{u}_N \in \mathbf{V}_N$:

$$\mathbf{u}_N = \mathbf{u}_{\text{ref}} + \sum_{k=1}^N \hat{u}_k(\boldsymbol{\mu}) \boldsymbol{\zeta}^k = \mathbf{u}_{\text{ref}} + \mathbf{V}_N \hat{\mathbf{u}}. \quad (1.6)$$

If $\mathbf{u}_{n_d} \in \mathbf{X}_\delta$, then its approximation is :

$$\mathbf{u}_{n_d} \simeq \mathbf{u}_{\text{ref}} + \mathbf{V}_N \hat{\mathbf{u}}.$$

We define the *model error* (or *bias*) of the ROM as

$$\boldsymbol{\eta} = \mathbf{u}_{n_d} - \mathbf{u}_N \in \mathbf{V}_N^\perp, \quad (1.22)$$

such that

$$\mathbf{u}_{n_d}(\boldsymbol{\mu}) = \mathbf{u}_N(\boldsymbol{\mu}) + \boldsymbol{\eta}(\boldsymbol{\mu}) = \mathbf{u}_{\text{ref}} + \mathbf{V}_N \hat{\mathbf{u}}(\boldsymbol{\mu}) + \boldsymbol{\eta}(\boldsymbol{\mu}). \quad (1.23)$$

The goal of what follows is to find a way to estimate the error $\boldsymbol{\eta}$. Toward that goal, we'll make use of the Talyor expansion [AGH+26].

Let $\boldsymbol{\mu} \in \mathbb{P}$, let $u(\boldsymbol{\mu}) \in X$ be the solution of Eq.(1.1). Assuming the map $\boldsymbol{\mu} \in \mathbb{P} \mapsto u(\cdot, \cdot; \boldsymbol{\mu})$ is smooth enough, we can apply a 1st order Taylor expansion with respect to $\boldsymbol{\mu}$:

$$u(\boldsymbol{\mu}) \underset{\|\boldsymbol{\mu} - \boldsymbol{\mu}_0\|_{\mathbb{R}^P} \rightarrow 0}{=} u(\boldsymbol{\mu}_0) + \nabla_{\boldsymbol{\mu}} u(\boldsymbol{\mu}_0) \cdot (\boldsymbol{\mu} - \boldsymbol{\mu}_0) + o(\|\boldsymbol{\mu} - \boldsymbol{\mu}_0\|_{\mathbb{R}^P}). \quad (1.24)$$

where $\nabla_{\mu}(\cdot)$ is the differentiation with respect to the parameter dependence.

We get the following analogy between the approximation Eq.(1.6) and Eq.(1.24) :

$$\mathbf{u}_{\text{ref}} \equiv \mathbf{u}(\mu_0), \quad (1.25)$$

$$V_N \hat{\mathbf{u}}(\mu) \equiv \nabla_{\mu}(u)(\mu_0) \cdot (\mu - \mu_0), \quad (1.26)$$

$$\boldsymbol{\eta}(\mu) \equiv o(\|\mu - \mu_0\|_{\mathbb{R}^P}) \quad \text{as } \|\mu - \mu_0\|_{\mathbb{R}^P} \rightarrow 0 \quad (1.27)$$

This is proven to be true locally [AGH+26]. Since $\nabla_{\mu}(u)(\mu_0) \in \mathcal{L}(\mathbb{R}^P, X)$, we expect a reduced basis approximating $u(\mu)$ for $\mu \simeq \mu_0$ to need $N \simeq P$ vectors to reconstruct, at least locally around μ_0 , the truth solution.

Hence it is wise to consider $N \gtrsim P$ modes for a reduced basis to approximate $u(\mu)$ if $\mu \simeq \mu_0$.

It is possible to go further with Taylor's expansion to get a better intuition of what $\boldsymbol{\eta}$ can be :

$$u(\mu) \underset{\|\mu - \mu_0\|_{\mathbb{R}^P} \rightarrow 0}{=} u(\mu_0) + \nabla_{\mu}(u)(\mu_0) \cdot (\mu - \mu_0) + \frac{1}{2} \nabla_{\mu}^2(u)(\mu_0)(\mu - \mu_0) \otimes (\mu - \mu_0) + o(\|\mu - \mu_0\|_{\mathbb{R}^P}^2). \quad (1.28)$$

Thus, we can expect for $\mu \simeq \mu_0$:

$$\boldsymbol{\eta}(\mu) = \frac{1}{2} \nabla_{\mu}^2(u)(\mu_0)(\mu - \mu_0) \otimes (\mu - \mu_0) + o(\|\mu - \mu_0\|_{\mathbb{R}^P}^2) \simeq \frac{1}{2} \nabla_{\mu}^2(u)(\mu_0)(\mu - \mu_0) \otimes (\mu - \mu_0) \quad (1.29)$$

If $u \in X$ with X an Hilbert, then we can write for $(\zeta^i)_i \in X^{\mathbb{N}}$:

$$u(\mu) = \sum_{i \in \mathbb{N}} \hat{u}_i(\mu) \zeta^i.$$

We can also write in the discrete case :

$$\mathbf{u}_{n_d}(\mu) = \mathbf{u}_{\text{ref}} + \sum_{1 \leq i \leq n_d} \hat{u}_i(\mu)(\mu) \zeta^i.$$

Thus we can write the equivalence :

$$\boldsymbol{\eta}(\mu) = \frac{1}{2} \nabla_{\mu}^2(u)(\mu_0)(\mu - \mu_0) \otimes (\mu - \mu_0) + o(\|\mu - \mu_0\|_{\mathbb{R}^P}^2) \underset{\|\mu - \mu_0\|_{\mathbb{R}^P} \rightarrow 0}{=} \sum_{i > P} \hat{u}_i(\mu) \zeta^i. \quad (1.30)$$

We observe that the coefficients appearing in $\boldsymbol{\eta}(\mu)$ are mixed term of the P firsts coefficients. The equivalence with the reduced basis coefficients suggests the same dependence of the last coefficients from the first ones. This is proven in [AGH+26] :

Theorem 1.7: True manifold dimension of the coefficients space

For $\mu_0 \in \mathbb{P} \subset \mathbb{R}^P$, there exist $r > 0$ such that for any $\mu \in \mathcal{B}_{\|\cdot\|_2}(\mu_0, r)$, the coefficients $\hat{u}_i(\mu)$, $i > n$ are function of $\hat{u}_1(\mu_0), \dots, \hat{u}_P(\mu_0)$, i.e. for any $i \in \mathbb{N}$, there exists $h_i : \mathbb{R}^P \rightarrow \mathbb{R}$ such that $\hat{u}_i = h_i(\hat{u}_1, \dots, \hat{u}_P)$.

We consider the following conjecture :

There exists $n \simeq P$ such that the coefficients $\hat{u}_i$, $i > n_0$ are function of $\hat{u}_1, \dots, \hat{u}_{n_0}$, i.e. for any $N \in \mathbb{N}$, there exists $h_n : \mathbb{R}^{n_0} \rightarrow \mathbb{R}$ such that $\hat{u}_n = h_n(\hat{u}_1, \dots, \hat{u}_{n_0})$.

Some efforts have already been made in that direction, e.g. using a quadratic approximation [SP24; JAI+17] or using encoding-decoding framework [BNS25; ARE+26]. Following [ARE+26], our goal is to approximate h_n with Machine Learning methods.

2) Decomposition in linear and nonlinear basis

Choose $n \in \mathbb{N}^*$ and $\bar{n} \in \mathbb{N}^*$ such that $N = n + \bar{n}$ is big enough to have :

$$\|\mathbf{S} - \mathbf{S}_N\|_F \leq \epsilon,$$

where $\mathbf{S}$ and $\mathbf{S}_N$ are defined by Theorem 1.6. Recall that $\hat{\mathbf{u}} = \mathbf{V}_N^T \mathbf{u} \in \mathbb{R}^N$ was defined as the N coefficients in the POD basis.

We define for $1 \leq i \leq n$ and $1 \leq j \leq \bar{n}$:

$$\alpha_i = \hat{u}_i, \quad \beta_j = \hat{u}_{j+N}.$$

The hypothesis of Theorem 1.7 suggests that there exists $n_0 \simeq P$ such that for any $n \geq n_0$ we can construct h which depends on n with

$$\boldsymbol{\beta} = h(\boldsymbol{\alpha}) \in \mathbb{R}^{\bar{n}},$$

where

$$\boldsymbol{\alpha} = \begin{bmatrix} \alpha_1 \\ \vdots \\ \alpha_n \end{bmatrix} \quad \text{and} \quad \boldsymbol{\beta} = \begin{bmatrix} \beta_1 \\ \vdots \\ \beta_{\bar{n}} \end{bmatrix}.$$

Suppose knowing only the first $n \simeq P$ coefficients $(\alpha_i)_{1 \leq i \leq n}$ is enough to reconstruct the following $\bar{n}$ coefficients $(\beta_i)_{1 \leq i \leq \bar{n}}$.

To highlights this separation of the coefficients, we consider two parts for the POD basis :

- The linear part $\mathbf{V} = [\boldsymbol{\zeta}^1 \quad \dots \quad \boldsymbol{\zeta}^n]$ consists of the n first columns of $\mathbf{V}_N$.
- The nonlinear part $\bar{\mathbf{V}} = [\boldsymbol{\zeta}^{n+1} \quad \dots \quad \boldsymbol{\zeta}^{\bar{n}}]$ consists of the columns $n+1$ to $N = \bar{n} + n$ of $\mathbf{V}_N$.

Thus if $\mathbf{u}_N - \mathbf{u}_{\text{ref}} \in \mathbf{V}_N =: [\mathbf{V} \quad \bar{\mathbf{V}}]$;

$$\mathbf{u}_N = \mathbf{u}_{\text{ref}} + \mathbf{V}\boldsymbol{\alpha} + \bar{\mathbf{V}}\boldsymbol{\beta} = \mathbf{u}_{\text{ref}} + \underbrace{\mathbf{V}\boldsymbol{\alpha}}_{\text{linear coefficients}} + \underbrace{\bar{\mathbf{V}}h(\boldsymbol{\alpha})}_{\text{nonlinear in } \boldsymbol{\alpha}}.$$

From now on, $\boldsymbol{\alpha}$ will always denote the linear coefficients in front on the linear basis $\mathbf{V}$, and $\boldsymbol{\beta}$ will denote the coefficients that can be represented by a function $h(\boldsymbol{\alpha})$ in front of $\bar{\mathbf{V}}$.

The error between u_{n_d} and this u_N is as follow :

$$\begin{aligned} \|\mathbf{u}_{n_d} - \mathbf{u}_N\|_X^2 &= \left\| \underbrace{\mathbf{u}_{n_d} - \mathbf{u}_{\text{POD}}}_{\in \mathbf{V}_N^\perp} + \underbrace{\mathbf{u}_{\text{POD}} - \mathbf{u}_N}_{\in \mathbf{V}_N} \right\|_X^2 \\ &= \|\mathbf{u}_{n_d} - \mathbf{u}_{\text{POD}}\|_X^2 + \|\mathbf{V}\boldsymbol{\alpha} + \bar{\mathbf{V}}\boldsymbol{\beta} - \mathbf{V}\boldsymbol{\alpha} - \bar{\mathbf{V}}h(\boldsymbol{\alpha})\|_X^2 \\ &= \|\boldsymbol{\eta}_{\text{POD}}\|_X^2 + \|\bar{\mathbf{V}}(\boldsymbol{\beta} - h(\boldsymbol{\alpha}))\|_X^2 \\ &= \|\boldsymbol{\eta}_{\text{POD}}\|_X^2 + \|\boldsymbol{\beta} - h(\boldsymbol{\alpha})\|_2^2. \end{aligned}$$

Thus the learning error will be controlled by the model error of the POD $\boldsymbol{\eta}_{\text{POD}}$ and the learning error between $h(\boldsymbol{\alpha})$ and $\boldsymbol{\beta}$:

$$\|\mathbf{u}_{n_d} - \mathbf{u}_N\|_X^2 = \|\boldsymbol{\eta}_{\text{POD}}\|_X^2 + \|\boldsymbol{\beta} - h(\boldsymbol{\alpha})\|_2^2. \quad (1.31)$$

Part 2

Nonlinear compressed reduced basis method for EIM

I - Sensor-observed system

1) Observation function

We actually do not have access to the continuous solution $u \in X$ but only of observations of the system. We set $m \in \mathbb{N}^*$ the number of sensors used to observe the system and assume that we directly observe u° ; i.e. there exists a function

$$H : u \in X \rightarrow \mathbb{R}^m,$$

called *observation function* such that any observation $\mathbf{y}^\circ \in \mathbb{R}^m$ of the system can be described as

$$\mathbf{y}^\circ = H(u^\circ)^8$$

for a certain $u^\circ \in \mathcal{M}$. In what follow, we suppose $H|_{\mathbb{R}^{n_d}} = \mathbf{P}$ is well-defined and linear. In our case, we consider that H is a linear pointwise observation such that for $\mathbf{u} \in \mathbb{R}^{n_d}$:

$$H(\mathbf{u}) := \mathbf{P}\mathbf{u} := \begin{bmatrix} u(\mathbf{x}_{i_1}) \\ \vdots \\ u(\mathbf{x}_{i_m}) \end{bmatrix} = \begin{bmatrix} u_{i_1} \\ \vdots \\ u_{i_m} \end{bmatrix}.$$

Where for $x_0 \leq \mathbf{x}_{i_1} \leq \dots \leq \mathbf{x}_{i_m} \leq \mathbf{x}_{n_d}$, $\mathbf{x}_{i_j}$ is the location of the j -th sensor in the discretized space. $\mathbf{P} : \mathbb{R}^{n_d} \rightarrow \mathbb{R}^m$ is a $m \times n_d$ matrix that selects the coefficients $(i_k)_{1 \leq k \leq m}$ of a vector in $\mathbb{R}^{n_d}$. To account for measurements errors, we define the noisy observation

$$\mathbf{y}_\varepsilon^\circ := \mathbf{y}^\circ + \mathbf{e} \in \mathbb{R}^m. \quad (2.1)$$

The error is supposed to be **relative error** modeled as a gaussian noise of scale $\sigma > 0$ that we assume to be centered and independent between each sensors. At sensor location i , the error is

$$\mathbf{e}_i = (\mathbf{y}^\circ)_i \cdot \varepsilon_i, \quad \varepsilon_i \sim \mathcal{N}(0, \sigma^2).$$

Thus the noisy observation at location i is :

$$(\mathbf{y}_\varepsilon^\circ)_i = (\mathbf{y}^\circ)_i + (\mathbf{y}^\circ)_i \cdot \varepsilon_i = (\mathbf{P}\mathbf{u}^\circ)_i + \mathbf{e}_i. \quad (2.2)$$

We can define the error vector as a gaussian vector :

$$\boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}_{R^m}, \sigma^2 \mathbf{I}_m),$$

This can be rewritten with the Hadamard product :

$$\mathbf{y}_\varepsilon^\circ = \mathbf{y}^\circ + \mathbf{y}^\circ \odot \boldsymbol{\varepsilon} = \mathbf{P}\mathbf{u}^\circ + \mathbf{e}. \quad (2.3)$$

where

$$\mathbf{e} := \mathbf{P}\mathbf{u}^\circ \odot \boldsymbol{\varepsilon}, \quad \boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}_{R^m}, \sigma^2 \mathbf{I}_m).$$

From that definition we have the relative error :

$$\mathbb{E} \left[\frac{\|\mathbf{e}\|_2^2}{\|\mathbf{y}^\circ\|_2^2} \right] = \sigma^2. \quad (2.4)$$

⁸Note we suppose here that u models exactly the real observed system and admits no modeling error.

2) Field reconstruction from partial observations

The EIM [BAR+04; MM13] propose a greedy selection of the sensors locations and an optimal reconstruction of the observed values.

a) Empirical Interpolation Method

The original EIM [BAR+04] method identifies iteratively couples of parameters and (magic) points : $(\boldsymbol{\mu}^1, \mathbf{x}^1), (\boldsymbol{\mu}^2, \mathbf{x}^2), \dots, (\boldsymbol{\mu}_{\text{EIM}}^N, \mathbf{x}_{\text{EIM}}^N)$ in a recursive manner as follows :

Assume $\{(\boldsymbol{\mu}^i, \mathbf{x}^i)\}_{i=1, \dots, n}$ are determined, and for any $\boldsymbol{\mu}$ we can define the interpolation operator $\mathcal{I}_n$ such that

$$\mathcal{I}_n[u(\cdot; \boldsymbol{\mu})] = \sum_{i=1}^n \gamma_i u(\cdot; \boldsymbol{\mu}^i) \quad (2.5)$$

with

$$\forall k = 1, \dots, n, \quad \mathcal{I}_n[u(\cdot; \boldsymbol{\mu})](\mathbf{x}^k) = u(\cdot; \boldsymbol{\mu})(\mathbf{x}^k) \quad (2.6)$$

then $\boldsymbol{\mu}^{n+1}$ is defined as

$$\boldsymbol{\mu}^{n+1} = \operatorname{argmax}_{\boldsymbol{\mu}} \|u(\cdot; \boldsymbol{\mu}) - \mathcal{I}_n[u(\cdot; \boldsymbol{\mu})]\|_X \quad (2.7)$$

and

$$\mathbf{x}^{n+1} = \operatorname{argmax}_{\mathbf{x}} |u(\mathbf{x}; \boldsymbol{\mu}^{n+1}) - \mathcal{I}_n[u(\cdot; \boldsymbol{\mu}^{n+1})](\mathbf{x})|. \quad (2.8)$$

Remark :

Note that it is interesting to introduce the basis

$$\xi^i = \frac{u(\cdot; \boldsymbol{\mu}^i) - \mathcal{I}_{i-1}[u(\cdot; \boldsymbol{\mu}^i)]}{u(\mathbf{x}^i; \boldsymbol{\mu}^i) - \mathcal{I}_{i-1}[u(\cdot; \boldsymbol{\mu}^i)](\mathbf{x}^i)} \quad (2.9)$$

since then, for any $\boldsymbol{\mu}$,

$$\mathcal{I}_n[u(\cdot; \boldsymbol{\mu})] = \mathcal{I}_{n-1}[u(\cdot; \boldsymbol{\mu})] + [u(\mathbf{x}^n; \boldsymbol{\mu}) - \mathcal{I}_{n-1}[u(\cdot; \boldsymbol{\mu})](\mathbf{x}^n)] \xi^n. \quad (2.10)$$

Aside from this original EIM method, was also proposed an alternative approach when an ordered appropriate reduced basis $\{\zeta^i\}_{i=1}^N$ is given. There is no point in identifying the $\boldsymbol{\mu}^j$, and only the "magic point" are determined through the following recursive method :

$$\mathbf{x}^1 = \operatorname{argmax}_{\mathbf{x}} |\zeta^1(\mathbf{x})| \quad (2.11)$$

then

$$\xi^1 = \frac{\zeta^1}{\zeta^1(\mathbf{x}^1)}$$

and

$$\mathcal{I}_1[u(\cdot; \boldsymbol{\mu})] = u(\mathbf{x}^1; \boldsymbol{\mu}) \xi^1. \quad (2.12)$$

Next, assuming that $(\xi^1, \mathbf{x}^1), (\xi^2, \mathbf{x}^2), \dots, (\xi^n, \mathbf{x}^n)$ have been built and, for any $\boldsymbol{\mu}$:

$$\mathcal{I}_n[u(\cdot; \boldsymbol{\mu})] = \mathcal{I}_{n-1}[u(\cdot; \boldsymbol{\mu})] + [u(\mathbf{x}^n; \boldsymbol{\mu}) - \mathcal{I}_{n-1}[u(\cdot; \boldsymbol{\mu})](\mathbf{x}^n)] \xi^n. \quad (2.13)$$

Hence

$$\mathcal{I}_n[u(\cdot; \boldsymbol{\mu})] = \mathcal{I}_1[u(\cdot; \boldsymbol{\mu})] + \sum_{i=1}^n [u(\mathbf{x}^i; \boldsymbol{\mu}) - \mathcal{I}_{i-1}[u(\cdot; \boldsymbol{\mu})](\mathbf{x}^i)] \xi^i. \quad (2.14)$$

We construct

$$\mathbf{x}^{n+1} = \operatorname{argmax}_{\mathbf{x}} |\zeta^{n+1}(\mathbf{x}) - \mathcal{I}_n[\zeta^{n+1}](\mathbf{x})| \quad (2.15)$$

and

$$\xi^{n+1} = \frac{\zeta^{n+1} - \mathcal{I}_n[\zeta^{n+1}]}{\zeta^{n+1}(\mathbf{x}^{n+1}) - \mathcal{I}_n[\zeta^{n+1}](\mathbf{x}^{n+1})} \quad (2.16)$$

b) General Empirical Interpolation Method

The original General Empirical Interpolation method (GEIM) [MM13] identifies iteratively couples of parameters and linear forms : $(\boldsymbol{\mu}^1, \boldsymbol{\sigma}^1), (\boldsymbol{\mu}^2, \boldsymbol{\sigma}^2), \dots, (\boldsymbol{\mu}_{\text{EIM}}^N, \boldsymbol{\sigma}_{\text{EIM}}^N)$ in a recursive manner as follows : Assume $\{(\boldsymbol{\mu}^i, \boldsymbol{\sigma}^i)\}_{i=1, \dots, n}$ are determined, and for any $\boldsymbol{\mu}$ we can define the interpolation operator $\mathcal{J}_n$ such that

$$\mathcal{J}_n[u(\cdot; \boldsymbol{\mu})] = \sum_{i=1}^n \gamma_i u(\cdot; \boldsymbol{\mu}^i) \quad (2.17)$$

with

$$\forall k = 1, \dots, n, \quad \boldsymbol{\sigma}^k [\mathcal{J}_n[u(\cdot; \boldsymbol{\mu})]] = \boldsymbol{\sigma}^k [u(\cdot; \boldsymbol{\mu})]. \quad (2.18)$$

Then $\boldsymbol{\mu}^{n+1}$ is defined as

$$\boldsymbol{\mu}^{n+1} = \operatorname{argmax}_{\boldsymbol{\mu}} \|u(\cdot; \boldsymbol{\mu}) - \mathcal{J}_n[u(\cdot; \boldsymbol{\mu})]\|_X, \quad (2.19)$$

and

$$\boldsymbol{\sigma}^{n+1} = \operatorname{argmax}_{\boldsymbol{\sigma}} |\boldsymbol{\sigma} [u(\cdot; \boldsymbol{\mu}^{n+1}) - \mathcal{J}_n[u(\cdot; \boldsymbol{\mu}^{n+1})]]|. \quad (2.20)$$

Remark :

Note that it is interesting to introduce the basis

$$\xi^i = \frac{u(\cdot; \boldsymbol{\mu}^i) - \mathcal{J}_{i-1}[u(\cdot; \boldsymbol{\mu}^i)]}{\boldsymbol{\sigma}^i [u(\boldsymbol{x}^i; \boldsymbol{\mu}^i) - \mathcal{J}_{i-1}[u(\cdot; \boldsymbol{\mu}^i)]]} \quad (2.21)$$

since then, for any $\boldsymbol{\mu}$,

$$\mathcal{J}_n[u(\cdot; \boldsymbol{\mu})] = \mathcal{J}_{n-1}[u(\cdot; \boldsymbol{\mu})] + \boldsymbol{\sigma}^n [u(\boldsymbol{x}^n; \boldsymbol{\mu}) - \mathcal{J}_{n-1}[u(\cdot; \boldsymbol{\mu})]] \xi^n \quad (2.22)$$

Aside from this original GEIM method, was also proposed an alternative approach when an ordered appropriate reduced basis $\{\zeta^i\}_{i=1}^N$ is given. There is no point in identifying the $\boldsymbol{\mu}^j$, and only the "magic moments" are determined through the following recursive method

$$\boldsymbol{\sigma}^1 = \operatorname{argmax}_{\boldsymbol{\sigma}} |\boldsymbol{\sigma} [\zeta^1]| \quad (2.23)$$

then

$$\xi^1 = \frac{\zeta^1}{\boldsymbol{\sigma}^1 [\zeta^1]}$$

and

$$\mathcal{J}_1[u(\cdot; \boldsymbol{\mu})] = \boldsymbol{\sigma}^1 u(\boldsymbol{x}^1; \boldsymbol{\mu}) \xi^1. \quad (2.24)$$

Next, assuming that $(\xi^1, \boldsymbol{x}^1), (\xi^2, \boldsymbol{x}^2), \dots, (\xi^n, \boldsymbol{x}^n)$ have been built and, for any $\boldsymbol{\mu}$,

$$\mathcal{J}_n[u(\cdot; \boldsymbol{\mu})] = \mathcal{J}_{n-1}[u(\cdot; \boldsymbol{\mu})] + \boldsymbol{\sigma}^n [u(\cdot; \boldsymbol{\mu}) - \mathcal{J}_{n-1}[u(\cdot; \boldsymbol{\mu})]] \xi^n. \quad (2.25)$$

Hence

$$\mathcal{J}_n[u(\cdot; \boldsymbol{\mu})] = \mathcal{J}_1[u(\cdot; \boldsymbol{\mu})] + \sum_{i=1}^n \boldsymbol{\sigma}^i [u(\cdot; \boldsymbol{\mu}) - \mathcal{J}_{i-1}[u(\cdot; \boldsymbol{\mu})]] \xi^i. \quad (2.26)$$

We construct

$$\boldsymbol{\sigma}^{n+1} = \operatorname{argmax}_{\boldsymbol{\sigma}} |\boldsymbol{\sigma} [\zeta^{n+1} - \mathcal{J}_n[\zeta^{n+1}]]| \quad (2.27)$$

and

$$\xi^{n+1} = \frac{\zeta^{n+1} - \mathcal{J}_n[\zeta^{n+1}]}{\boldsymbol{\sigma}^{n+1} [\zeta^{n+1} - \mathcal{J}_n[\zeta^{n+1}]]}. \quad (2.28)$$

Following the *gappy EIM* [PD20] we construct $\boldsymbol{P} : \mathbb{R}^{n_d} \rightarrow \mathbb{R}^m$, $m \geq N_{\text{EIM}}$ using the EIM to add more observations than there are basis functions.

3) Succinct analysis of the noisy EIM error

Let n_d be the discrete spacetime dimension of the full-order problem, and let X be the associated state space equipped with the norm $\|\cdot\|_X$.

We recall that the POD basis $(\xi_i)_{i=1}^{n_d}$ fully captures the discrete vector $\mathbf{u}^\circ \in \mathbb{R}^{n_d}$.

Let $N_{\text{EIM}} \geq 1$ be the number of EIM functions $(\psi^i)_{i=1}^{N_{\text{EIM}}}$ who are the columns of $\mathbf{V}_{\text{EIM}} \in \mathbb{R}^{n_d \times N_{\text{EIM}}}$.

A function $\mathbf{u}^\circ \in \mathbb{R}^{n_d}$ deviated from a known reference state $\mathbf{u}_{\text{ref}} \in \mathbb{R}^{n_d}$ can be described by :

$$\mathbf{u}^\circ - \mathbf{u}_{\text{ref}} = \sum_{i=1}^{N_{\text{EIM}}} \gamma_i^\circ \xi^i + \sum_{i=N_{\text{EIM}}+1}^{n_d} \alpha_i^\circ \xi^i = \mathbf{V}_{\text{EIM}} \boldsymbol{\gamma}^\circ + \boldsymbol{\eta},$$

where $\boldsymbol{\eta} = \sum_{i=N_{\text{EIM}}+1}^{n_d} \alpha_i^\circ \xi^i$ represents the model error such that $\boldsymbol{\eta} \perp \mathbf{V}_{\text{EIM}}$.

Let $\mathbf{P} \in \mathbb{R}^{m \times n_d}$ be the sensor placement matrix sampling $m \geq n_{\text{EIM}}$ locations. The noisy sensor observation vector $\mathbf{y}_\varepsilon^\circ \in \mathbb{R}^m$ is modeled as :

$$\mathbf{y}_\varepsilon^\circ = \mathbf{P} \mathbf{u}^\circ + \mathbf{e},$$

where $\mathbf{e} \in \mathbb{R}^m$ is a relative measurement noise vector defined by $\mathbf{e} = \sigma(\mathbf{P} \mathbf{u}^\circ) \odot \boldsymbol{\epsilon}$, with noise level $\sigma > 0$ and $\boldsymbol{\epsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_m)$.

The gappy EIM reconstruction operator $\mathbf{M} \in \mathbb{R}^{n_d \times m}$ is defined as $\mathbf{M} = \mathbf{V}_{\text{EIM}}(\mathbf{P} \mathbf{V}_{\text{EIM}})^\dagger$, which satisfies the fundamental interpolation identity $\mathbf{M} \mathbf{P} \mathbf{V}_{\text{EIM}} = \mathbf{V}_{\text{EIM}}$. The reconstructed field is given by :

$$\mathbf{u}_{\text{EIM}} = \mathbf{M}(\mathbf{y}_\varepsilon^\circ - \mathbf{P} \mathbf{u}_{\text{ref}}) + \mathbf{u}_{\text{ref}}.$$

By substituting the observation model into the reconstruction equation, we isolate the total approximation error:

$$\begin{aligned} \mathbf{u}^\circ - \mathbf{u}_{\text{EIM}} &= (\mathbf{u}^\circ - \mathbf{u}_{\text{ref}}) - \mathbf{M}(\mathbf{P}(\mathbf{u}^\circ - \mathbf{u}_{\text{ref}}) + \mathbf{e}) \\ &= (\mathbf{I} - \mathbf{M} \mathbf{P})(\mathbf{u}^\circ - \mathbf{u}_{\text{ref}}) - \mathbf{M} \mathbf{e} \\ &= (\mathbf{I} - \mathbf{M} \mathbf{P})(\mathbf{V}_{\text{EIM}} \boldsymbol{\gamma}^\circ + \boldsymbol{\eta}) - \mathbf{M} \mathbf{e}. \end{aligned}$$

Using the property $(\mathbf{I} - \mathbf{M} \mathbf{P}) \mathbf{V}_{\text{EIM}} = \mathbf{0}$, the projection error cleanly simplifies to :

$$\mathbf{u}^\circ - \mathbf{u}_{\text{EIM}} = (\mathbf{I} - \mathbf{M} \mathbf{P}) \boldsymbol{\eta} - \mathbf{M} \mathbf{e}.$$

Taking the total norm and applying the triangle inequality yields :

$$\|\mathbf{u}^\circ - \mathbf{u}_{\text{EIM}}\|_X \leq \|\mathbf{I} - \mathbf{M} \mathbf{P}\|_2 \|\boldsymbol{\eta}\|_X + \|\mathbf{M} \mathbf{e}\|_X.$$

Squaring both sides and taking the expectation with respect to the random noise field $\boldsymbol{\epsilon}$, we evaluate the scaling of the stochastic term :

$$\begin{aligned} \mathbb{E}_\epsilon \left[\|\mathbf{M} \mathbf{e}\|_X^2 \right] &= \mathbb{E}_\epsilon \left[\|\mathbf{M}(\mathbf{P} \mathbf{u}^\circ \odot \boldsymbol{\epsilon})\|_X^2 \right] \\ &\leq \|\mathbf{M}\|_2^2 \mathbb{E}_\epsilon \left[\|\mathbf{P} \mathbf{u}^\circ \odot \boldsymbol{\epsilon}\|_2^2 \right] \\ &= \|\mathbf{M}\|_2^2 \sum_{i=1}^m |(\mathbf{P} \mathbf{u}^\circ)_i|^2 \mathbb{E}_\epsilon [\epsilon_i^2] \\ &= \sigma^2 \|\mathbf{M}\|_2^2 \sum_{i=1}^m |(\mathbf{P} \mathbf{u}^\circ)_i|^2 \\ &\leq \sigma^2 \|\mathbf{M}\|_2^2 \|\mathbf{P}\|_2^2 \|\mathbf{u}^\circ\|_X^2. \end{aligned}$$

Normalizing by the energy of the true state $\|\mathbf{u}^\circ\|_X^2$ and using Young's inequality for products :

$$\frac{\mathbb{E}_\epsilon \left[\|\mathbf{u}^\circ - \mathbf{u}_{\text{EIM}}\|_X^2 \right]}{\|\mathbf{u}^\circ\|_X^2} \leq 2 \|\mathbf{I} - \mathbf{M}\mathbf{P}\|_2^2 \frac{\|\boldsymbol{\eta}\|_X^2}{\|\mathbf{u}^\circ\|_X^2} + 2\sigma^2 \|\mathbf{M}\|_2^2 \|\mathbf{P}\|_2^2. \quad (2.29)$$

Equation (2.29) highlights the fundamental EIM trade-off: increasing the number of modes N_{EIM} minimizes the relative truncation error $\|\boldsymbol{\eta}\|_X / \|\mathbf{u}^\circ\|_X$, but increases the operator norm bounds $\|\mathbf{M}\|_2$, thereby magnifying the sensor noise sensitivity proportional to σ^2 .

Bounds of $\|\mathbf{I} - \mathbf{M}\mathbf{P}\|_2$ and $\|\mathbf{M}\|_2$ exist with respect to the number of EIM basis functions N_{EIM} and the number of observations $m \geq N_{\text{EIM}}$ [PD20] .

4) Data assimilation methods

Figure 1 summarizes the workflow that will be described in the next part.

The goal is to build in an offline part the needed POD basis, the EIM basis and magic points and an encoder h trained via a Gaussian Process Regressor (GPR). Then, from a noisy partial observation of the system, we try different methods to recover the reduced POD basis coefficients

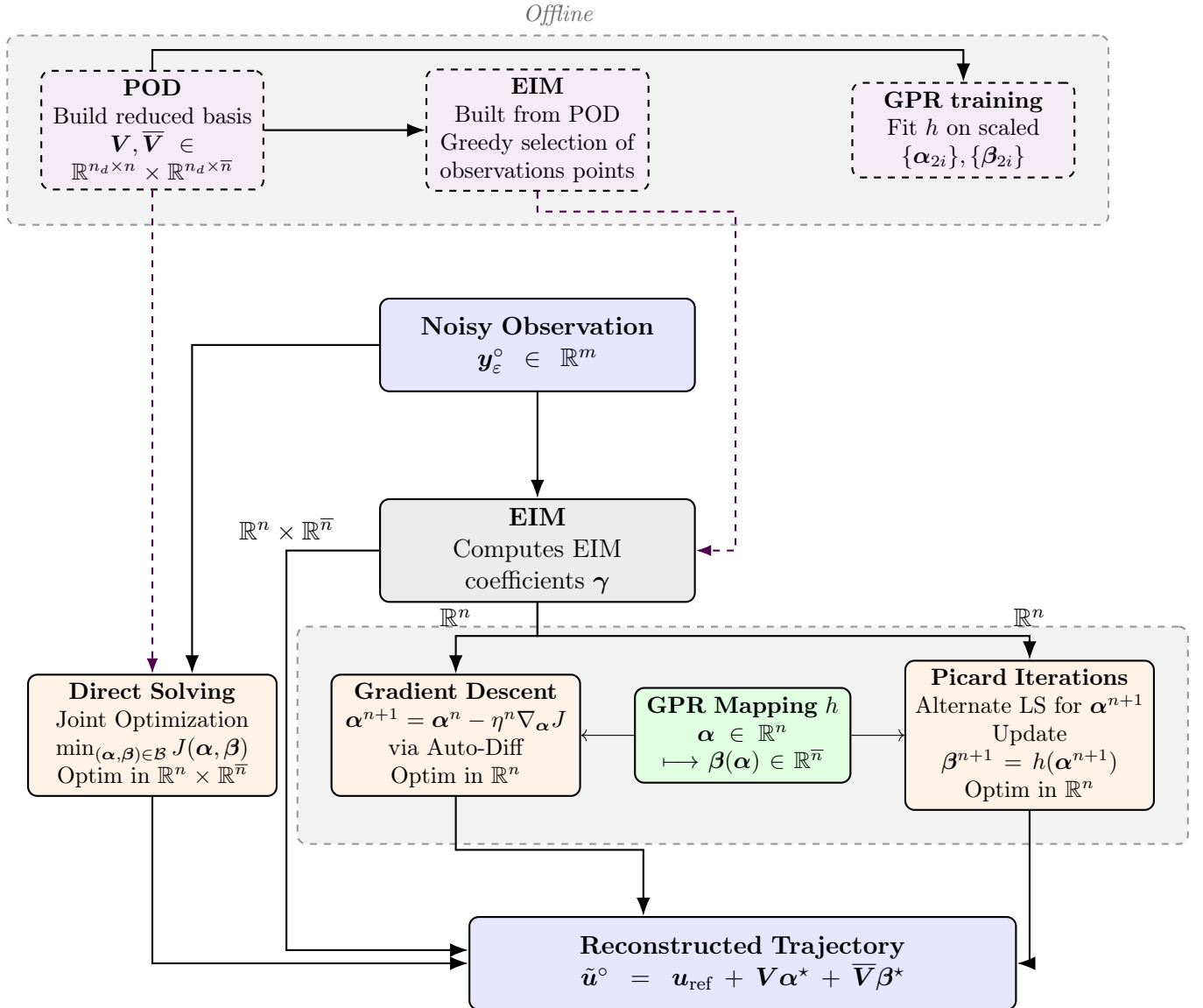


Figure 1: Architecture of the different proposed methods of assimilation.

We focus on retrieving the full trajectory $\mathbf{u}_{nd}$ from observations $\mathbf{y}_\varepsilon^\circ$ using the POD basis.

This has already been addressed with the Generalized Empirical Interpolation Method (EIM) [BAR+04; MM13; GON+19], GEIM with ML closure [NIU+25] and Parameterized Best Knowledge Model (PBDW) [MAD+15; GON+19] if accounting for modeling error.

Let $\mathbf{V}_N$ be a reduced basis. Recall that $\Pi_{\mathbf{V}_N}$ is the orthogonal projection onto $\mathbf{V}_N$.

Let $\mathbf{P} : \mathbb{R}^{nd} \rightarrow \mathbb{R}^m$ be a linear observation operator.

For an noisy observation $\mathbf{y}_\varepsilon^\circ$ find POD coefficients $(\boldsymbol{\alpha}^\circ, \boldsymbol{\beta}^\circ) \in \mathbb{R}^n \times \mathbb{R}^{\bar{n}}$ that are the POD coefficients of the true trajectory $\mathbf{u}^\circ \in \mathbb{R}^{nd}$:

$$\begin{cases} \mathbf{y}_\varepsilon^\circ &= \mathbf{P}\mathbf{u}^\circ + \mathbf{e}, \\ \tilde{\mathbf{u}}^\circ &= \mathbf{u}_{\text{ref}} + \Pi_{\mathbf{V}_N}(\mathbf{u}^\circ - \mathbf{u}_{\text{ref}}) = \mathbf{u}_{\text{ref}} + \mathbf{V}\boldsymbol{\alpha}^\circ + \bar{\mathbf{V}}\boldsymbol{\beta}^\circ. \end{cases} \quad (2.30)$$

Let $\boldsymbol{\alpha}_{\min}, \boldsymbol{\alpha}_{\max} \in \mathbb{R}^n$ and $\boldsymbol{\beta}_{\min}, \boldsymbol{\beta}_{\max} \in \mathbb{R}^{\bar{n}}$ be respectively the **componentwise** minimum and maximum values of $\boldsymbol{\alpha}, \boldsymbol{\beta}$ for all the snapshots in $\mathbb{P}_h$. Let $\delta_\alpha, \delta_\beta > 0$ be some box tolerance.

In what follow, we will use the following box constraint during optimizations [GON+19] :

$$\mathcal{B} := \prod_{i=1}^n [\alpha_{\min,i} - \delta_\alpha; \alpha_{\max,i} + \delta_\alpha] \times \prod_{i=1}^{\bar{n}} [\beta_{\min,i} - \delta_\beta; \beta_{\max,i} + \delta_\beta] \quad (2.31)$$

or

$$\tilde{\mathcal{B}} = \prod_{i=1}^n [\alpha_{\min,i} - \delta_\alpha; \alpha_{\max,i} + \delta_\alpha]. \quad (2.32)$$

When there is no ambiguity, we will use the notation $\mathcal{B}$ even when optimizing only on $\boldsymbol{\alpha}$.

We propose several methods that can be written in two forms for a cost function $J > 0$:

$$(a) \quad \text{Find } (\boldsymbol{\alpha}^*, \boldsymbol{\beta}^*) := \underset{(\boldsymbol{\alpha}, \boldsymbol{\beta}) \in \mathcal{B}}{\operatorname{argmin}} J(\boldsymbol{\alpha}, \boldsymbol{\beta}), \quad (2.33)$$

$$(b) \quad \text{Find } \boldsymbol{\alpha}^* := \underset{\boldsymbol{\alpha} \in \mathcal{B}}{\operatorname{argmin}} J(\boldsymbol{\alpha}), \quad \boldsymbol{\beta}^* = h(\boldsymbol{\alpha}^*). \quad (2.34)$$

a) Direct solving with full POD model

For an observed $\mathbf{y}_\varepsilon^\circ \in \mathbb{R}^m$, we solve the linear least square problem :

$$(\boldsymbol{\alpha}_{\text{POD}}^*, \boldsymbol{\beta}_{\text{POD}}^*) := \underset{(\boldsymbol{\alpha}, \boldsymbol{\beta}) \in \mathcal{B}}{\operatorname{argmin}} \frac{1}{2} \|\mathbf{y}_\varepsilon^\circ - \mathbf{P}(\mathbf{V}\boldsymbol{\alpha} + \bar{\mathbf{V}}\boldsymbol{\beta} + \mathbf{u}_{\text{ref}})\|_{\mathbf{L}^2(\mathbb{R}^m)}^2. \quad (2.35)$$

b) Gradient descent with ML model

We let $(\boldsymbol{\alpha}_0, \boldsymbol{\beta}_0) = \text{EIM}(\mathbf{y}_\varepsilon^\circ)$. Suppose the ML function h is differentiable, we apply a gradient descent to :

$$J(\boldsymbol{\alpha}) = \frac{1}{2} \|\mathbf{y}_\varepsilon^\circ - \mathbf{P}(\mathbf{V}\boldsymbol{\alpha} + \bar{\mathbf{V}}h(\boldsymbol{\alpha}) + \mathbf{u}_{\text{ref}})\|_{\mathbf{L}^2(\mathbb{R}^m)}^2. \quad (2.36)$$

$$\boldsymbol{\alpha}_{\text{GD}}^* := \underset{\boldsymbol{\alpha} \in \mathcal{B}}{\operatorname{argmin}} J(\boldsymbol{\alpha}). \quad (2.37)$$

The gradient with respect to the input $\boldsymbol{\alpha}$ is :

$$\nabla_{\boldsymbol{\alpha}} J = \left(\mathbf{P}\mathbf{V} + \mathbf{P}\bar{\mathbf{V}} \frac{\partial h}{\partial \boldsymbol{\alpha}} \right)^T (\mathbf{y}_\varepsilon^\circ - \mathbf{P}(\mathbf{V}\boldsymbol{\alpha} + \bar{\mathbf{V}}h(\boldsymbol{\alpha}) + \mathbf{u}_{\text{ref}})) \in \mathbb{R}^n. \quad (2.38)$$

The gradient $\frac{\partial h}{\partial \alpha} \in \mathbb{R}^{\bar{n} \times n}$ is done by automatic differentiation [NEI10; BAY+18] using the Tensorflow library⁹.

Let $\eta^n > 0$ be the learning rate defined in [BB88]. We update α^n as :

$$\alpha^{n+1} = \alpha^n - \eta^n \nabla_{\alpha} J(\alpha^n).$$

The process stops when $\|\nabla_{\alpha} J(\alpha^n)\| < \text{tol}$ or $n > n_{\max \text{ iter}}$ where tol and $n_{\max \text{ iter}}$ are arbitrary constants.

We denote the solution α_{GD}^* at the end of the process and set $\beta_{\text{GD}}^* = h(\alpha_{\text{GD}}^*)$ such that the final solution is given by $(\alpha_{\text{GD}}^*, \beta_{\text{GD}}^*)$.

c) Picard iterations

Let $\gamma_0 = \text{EIM}(\mathbf{y}_{\varepsilon}^{\circ})$. Let $\alpha_0 = \Pi_V(\mathbf{V}_{\text{EIM}}\gamma_0)$. We set $\beta_0 = h(\alpha_0)$.

Suppose we computed α_n, β_n . We solve the linear least square problem :

$$\alpha^{n+1} := \underset{\alpha \in \mathcal{B}}{\text{argmin}} \frac{1}{2} \left\| \mathbf{y}_{\varepsilon}^{\circ} - \mathbf{P}(\mathbf{V}\alpha + \bar{\mathbf{V}}\beta^n + \mathbf{u}_{\text{ref}}) \right\|_{\mathbf{L}^2(\mathbb{R}^m)}^2. \quad (2.39)$$

Once computed, we set $\beta^{n+1} := h(\alpha^{n+1})$.

The optimization process is repeated until $\frac{\|\alpha^{n+1} - \alpha^n\|_{\mathbb{R}^n}}{\|\alpha^n\|_{\mathbb{R}^n}} < \text{tol}$ or until the maximum number of iterations $n_{\max}$ has been reached.

We denote the solution at the end of the process $(\alpha_{\text{Picard}}^*, \beta_{\text{Picard}}^*) = (\alpha^{n_{\text{final}}+1}, \beta^{n_{\text{final}}+1})$ (or $\beta^{n_{\text{final}}}$).

5) Learning the mapping between coefficients

We use Gaussian Process Regression (GPR) [RW05] models to learn the mapping $\alpha \in \mathbb{R}^n \mapsto \beta(\alpha)$.

We use different kernels [RW05, Section 4.2] to learn the function :

- Constant kernel of scale ℓ :

$$\text{Cste}(\ell)(x_i, x_j) = \ell^2.$$

- Dot product kernel of scale ℓ :

$$\text{Dot}(\ell)(x_i, x_j) = \ell^2 + x_i x_j.$$

- Radial Basis Function kernel (RBF) with scale ℓ :

$$\text{RBF}(\ell)(x_i, x_j) = \exp(-|x_i - x_j'|^2 / 2\ell^2).$$

- Matern kernel with scale ℓ and regularity ν :

$$\text{Matern}(\nu)(x_i, x_j) = \frac{1}{\Gamma(\nu)2^{\nu-1}} \left(\frac{\sqrt{2\nu}}{\ell} |x_i - x_j| \right)^{\nu} K_{\nu}(x_i, x_j).$$

Here K_{ν} is a modified Bessel function, see [RW05, Chapter 4, Section 4.2] .

- White noise kernel of scale ϵ (here δ_0 is the Dirac operator at 0) :

$$\text{WhiteNoise}(\epsilon)(x_i, x_j) = \epsilon \delta_0(x_i, x_j).$$

⁹www.tensorflow.org/

The sum, dot product and power of kernels are allowed operations that are already implemented in the used libraries. The hyperparameters (scale, noise level, constants, ...) are automatically optimized within the fitting part. The GPR is trained on scaled input.

We use the scikit-learn¹⁰ gaussian process library for the training over the data.

If needed (e.g. for automatic differentiation), the sklearn GPR will be converted into tensorflow¹¹ using GPflow¹² GPR.

We suppose $\mu_1 < \mu_2 < \dots < \mu_{N_{\text{samples}}}$. The GPR is trained on $\mathbb{P}_{h,\text{train}} := \{\mu_{2i}\}_i$ then tested on $\mathbb{P}_{h,\text{test}} := \{\mu_{2i+1}\}_i$.

II - Numerical results

1) Test cases

a) Orthogonal modes, $P = 1, N = 2$

Let $I = [0; \pi] \subset \mathbb{R}$, $\mathbb{P} = [0; 1] \subset \mathbb{R}$ ($P = 1$). Here, the parameter space dimension is 1. Let $\lambda = 0.1$.

We consider a function u parameterized by $\mu \in \mathbb{P}$:

$$u(x; \mu) = \hat{u}_1(\mu)\Phi_1(x) + \hat{u}_2(\mu)\Phi_2(x) \quad (2.40)$$

$$= \sqrt{12} \left(\mu - \frac{1}{2} \right) \sqrt{\frac{2}{\pi}} \sin(x) + \lambda \sqrt{180} \left(\mu^2 - \mu + \frac{1}{6} \right) \sqrt{\frac{2}{\pi}} \cos(x). \quad (2.41)$$

Following Appendix B, the function u is its own exact POD approximation i.e. $\mathbf{u}_{n_d} = \mathbf{u}_n$ for $N = 2$ with modes :

$$\begin{cases} \Phi_1 &= \sqrt{\frac{2}{\pi}} \sin, \\ \Phi_2 &= \sqrt{\frac{2}{\pi}} \cos. \end{cases} \quad (2.42)$$

And coefficients :

$$\begin{cases} \alpha(\mu) := \hat{u}_1(\mu) &= \sqrt{12} \left(\mu - \frac{1}{2} \right), \\ \beta(\mu) := \hat{u}_2(\mu) &= \lambda \sqrt{180} \left(\mu^2 - \mu + \frac{1}{6} \right). \end{cases} \quad (2.43)$$

We discretize $[0; \pi]$ in $n_d = 4096$ points $x_0 = 0 \leq x_i = \frac{i}{4095} \leq x_{4095} = \pi$. The samples are $\mathbf{u}_{n_d} = [u(x_0) \dots u(x_{4095})]^T \in \mathbb{R}^{4096}$.

We construct $\mathbb{P}_h$ by uniformly sampling 2000 values in $[0; 1]$. We define $\mathbb{P}_{h,\text{train}} = \{\mu_{2i} / i \in \llbracket 0; 2000 \rrbracket\}$.

Up to sampling error, the discrete POD modes (Φ_1, Φ_2) will be :

$$\Phi_1^h = \pm \sqrt{2/\pi} \sin, \quad \Phi_2^h = \pm \sqrt{2/\pi} \cos.$$

We have the relation of dependance between α and β :

$$\beta(\alpha) = \lambda \frac{\sqrt{180}}{12} ((\alpha)^2 - 1). \quad (2.44)$$

Thus we will set $n = 1, \bar{n} = 1$ and $N = n + \bar{n} = 2$.

¹⁰<https://scikit-learn.org/>

¹¹www.tensorflow.org/

¹²[gpflow.github.io](https://github.com/gpflow/gpflow)

b) Advection of a gaussian bump, $P = 1$, $N = \infty$

Let $I = [-0.1; 1.1] = [a; b]$, $\mathbb{P} = [0; 1]$ ($P = 1$). Let $s > 0$.

We consider an accelerating advection of a gaussian bump :

$$u(x; \mu) = \frac{1}{(2\pi s^2)^{1/2}} \exp\left(-\frac{(x - \mu^2)^2}{2s^2}\right). \quad (2.45)$$

We choose $s^2 = 0.01$ such that $u(\mathbf{x} = -0.1; \mu = 0) = u(\mathbf{x} = 1.1; \mu = 1) \simeq 0$.

We discretize $[-0.1; 1.1]$ in $n_d = 4096$ points $x_0 = -0.1 \leq x_i = \frac{i}{4095} \leq x_{4095} = 1.1$. The samples are $\mathbf{u}_{n_d} = [u(x_0) \dots u(x_{4095})]^T \in \mathbb{R}^{4096}$.

We construct $\mathbb{P}_h$ by uniformly sampling 2000 values with the function $\mu \in [0; 1] \mapsto \mu^2$ which is equivalent to uniformly sampling $[0; 1]$ and consider $\exp(-(\mathbf{x} - \mu)^2)$. We define $\mathbb{P}_{h,\text{train}} = \{\mu_{2i} / i \in \llbracket 0; 2000 \rrbracket\}$.

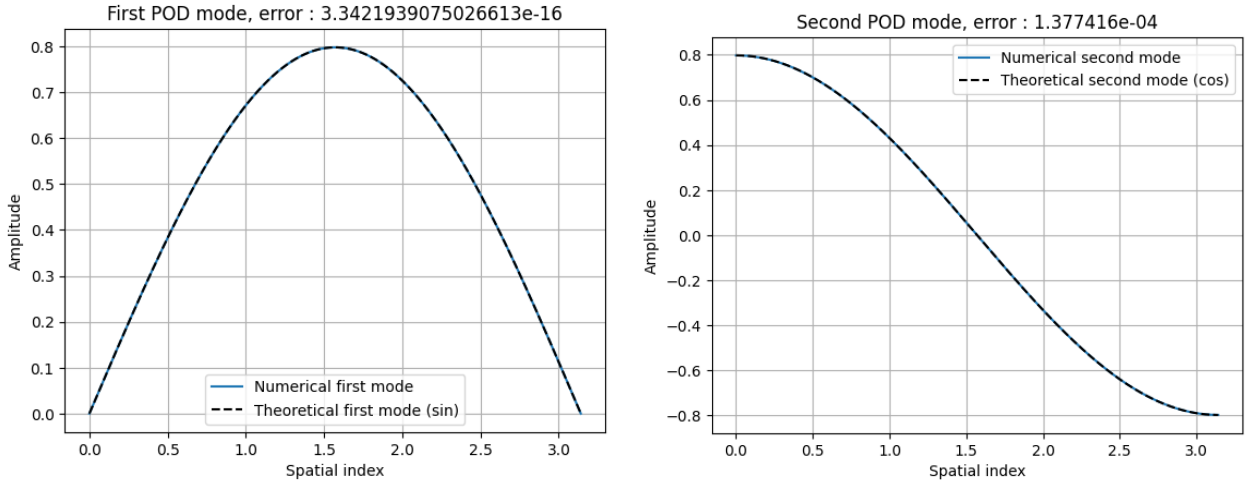
Appendix B shows that it is not possible to exactly recover $\mathcal{M} = \{u(\cdot; \mu) / \mu \in \mathbb{P}\}$ with a finite dimensional space, but it can be done up to any numerical precision.

2) Orthogonal modes results

a) POD and EIM analysis

We benchmark the POD implementation on the Orthogonal modes Eq.(2.41) for $N_{\text{train}} = 2000$.

Recall $\Phi_1 = \sqrt{\frac{2}{\pi}} \sin$ and $\Phi_2 = \sqrt{\frac{2}{\pi}} \cos$.



(a) First mode of the numerical POD Φ_1^h and the theoretical first mode Φ_1 .

(b) Second mode of the numerical POD Φ_2^h and the theoretical second mode Φ_2 .

Figure 2: Modes of the numerical POD and their $\|\cdot\|_X^2$ -error with respect to the theoretical mode.

The POD mode error is computed by $\|\Phi_i^h - \Phi_i\|_X$ following Eq.(3.3) of Appendix A.

The energy is $\mathcal{E}_1 = \sigma_1^2 = 1 - \lambda^2 = 0.99$ and $\mathcal{E}_2 = \sigma_2^2 = \lambda^2 = 0.01$, as stated in Appendix B :

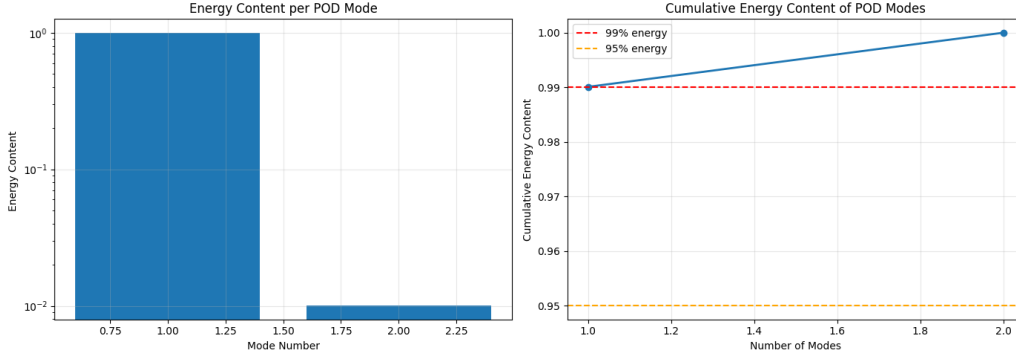


Figure 3: Energy histogram per modes (left). Cumulative energy $I = \sum_{i \leq n} \mathcal{E}_i / \sum_i \mathcal{E}_i$ (right).

We can verify Eq.(2.43) and Eq.(2.44) :

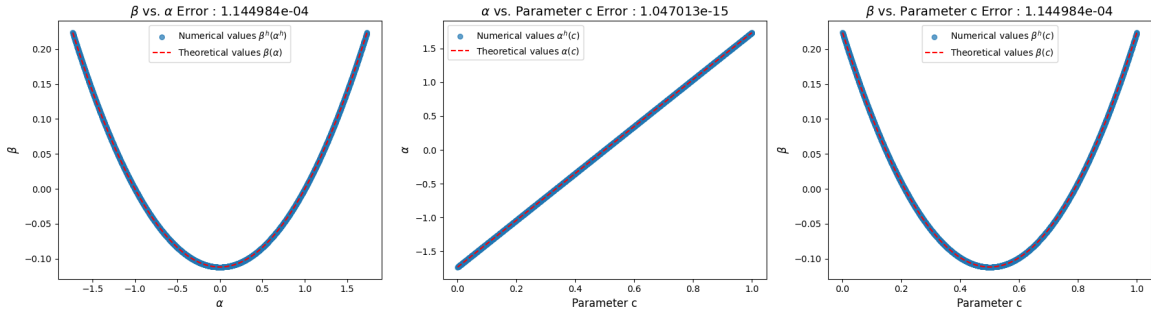


Figure 4: Comparison of the numerical coefficients with the theory of the orthogonal test case.

The error is computed as the Mean Relative Error Eq.(3.10) of the $\|\cdot\|_{L^2(\mathbb{P}_h)}$ norm with uniform sampling, see Appendix A.

Here, $c = \mu \in \mathbb{P}$ is the parameter, $N_s = N_{\text{train}} = 2000$, $\alpha = \alpha(c)$ and $\beta = \beta(c)$ are defined Eq.(2.43).

The error values are numerical precision on α and small enough on β to consider that the POD is well-implemented up to sampling error.

We check that the selected magic points of the EIM are $x_1 \in \frac{\pi}{2}$ and $x_2 \in \{0, \pi\}$ (Appendix B) :

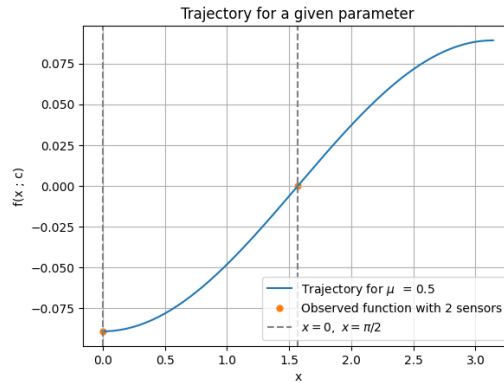


Figure 5: Trajectory $u^\circ(x; \mu)$ for a given $\mu = 0.5$, the observed function $\mathbf{y}^\circ$ and the sensors location.

We compare the results of EIM and POD in the case where there's no noise ($\sigma = 0$).

	Number of Modes	MRE	RMSE	Relative L^∞ error
POD	$n = 2$	7.78×10^{-16}	7.04×10^{-16}	2.20×10^{-14}
EIM	$m = n = 2$	1.99×10^{-14}	7.61×10^{-13}	2.19×10^{-14}

Table 2: Error metrics in the $\|\cdot\|_X$ norm.

Following Appendix A the MRE is done for the $\mathbf{L}^2(\mathbb{P}_h)$ -norm with $y = \mathbf{M}_\delta^{1/2} \mathbf{u}_{n_d}(\boldsymbol{\mu})$ and $v = \mathbf{M}_\delta^{1/2} \mathbf{u}_n(\boldsymbol{\mu})$. The RMSE is done for the same y, v . The relative L^∞ error is done on the space $\mathbf{L}^2(\mathbb{P}_h, \mathbf{L}^2(\Omega, \mathbb{R}^{n_d}))$.

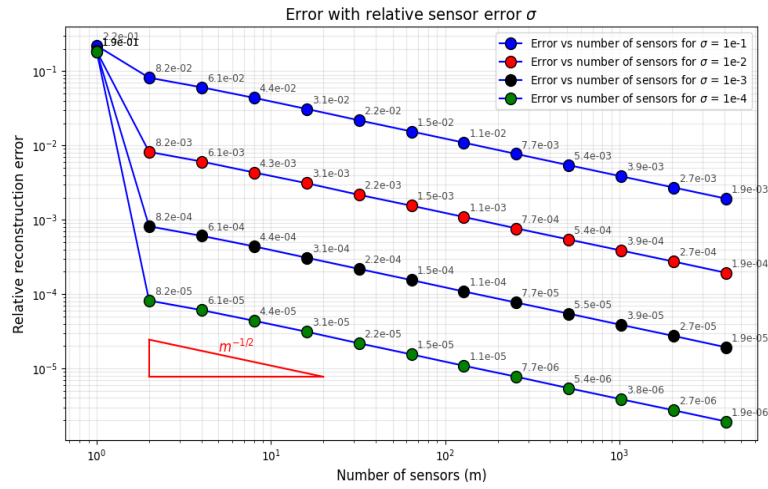


Figure 6: Relative error $\|\mathbf{u}^\circ - \mathbf{u}_{\text{EIM}}\|_X / \|\mathbf{u}^\circ\|_X$ over $\mathbb{P}_{h,\text{test}}$ for different value of the noise scale σ when the number of sensors m varies and $n_{\text{EIM}} = 2$ if $m > 1$, $n_{\text{EIM}} = 1$ if $m = 1$.

Note that the slope is $m^{-1/2}$ for each error scale and the error for $m = N = 2$ is of the scale of σ which is what the theory predict [PD20] and Eq.(2.29).

b) Learning the coefficients map

Recall we are trying to learn the mapping Eq.(2.44) :

$$\beta(\alpha) = \lambda \frac{\sqrt{180}}{12} ((\alpha)^2 - 1).$$

Kernel	$(n, \bar{n})$	MRE	RMSE	Relative L^∞ error
RBF + Cste**2	(1,1)	1.77×10^{-07}	5.14×10^{-09}	4.28×10^{-08}
Matern(nu=1.5) + Cste**2	(1,1)	4.22×10^{-08}	3.56×10^{-08}	4.89×10^{-06}
Matern(1, nu=5) + Cste**2 + White	(1,1)	2.85×10^{-09}	1.96×10^{-09}	1.66×10^{-07}
Dot ** 2	(1,1)	3.80×10^{-16}	5.22×10^{-16}	1.19×10^{-14}

Table 3: Error metrics of the GPR output $h(\alpha(\mu))$ with respect to the target value $\beta(\mu)$ for different kernels and $\mu \in \mathbb{P}_{h,\text{test}}$.

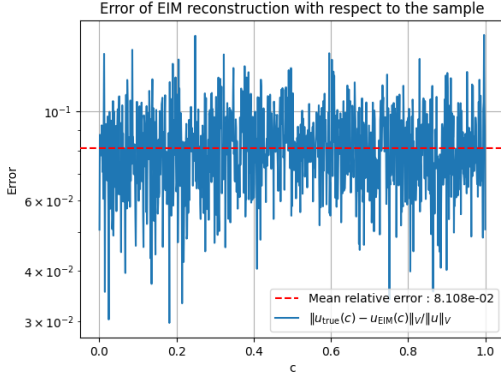
Note that DotProduct optimized its scale to $\ell = 0.999$ which implies the optimized kernel is a polynomial of the form $x^2 + \ell^2$ up to a scaling. That correspond almost exactly to $\alpha^2 + (-1)^2$.

With not much surprises, the polynomial kernel performs the best.

c) Online data assimilation

For each $\mu \in \mathbb{P}_{h,\text{test}}$ we generate $N_{\text{real}} = 10$ trajectories $\mathbf{u}_i^\circ(\mu)$ and a noisy observation $\mathbf{y}_i^\circ(\mu) = \mathbf{P}\mathbf{u}_i^\circ + \mathbf{e}_i$ with independent realization of the noise $\mathbf{e}$. We consider a relative scale $\sigma = 10^{-1}$. This construct $N_{\text{real}} \times n_{\text{test samples}} = 10^4$ noisy observations.

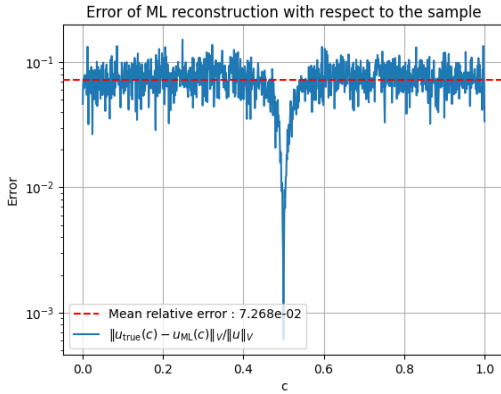
For $\mu \in \mathbb{P}$, $\epsilon_i(\mu) = \|\mathbf{u}_i^\circ(\mu) - \mathbf{u}_i^{\text{reconstructed}}\|_X / \|\mathbf{u}_i^\circ(\mu)\|_X$. Figure 7 plot $\epsilon(\mu) = \frac{1}{N_{\text{real}}} \sum_{i=1}^{N_{\text{real}}} \epsilon_i(\mu)$.



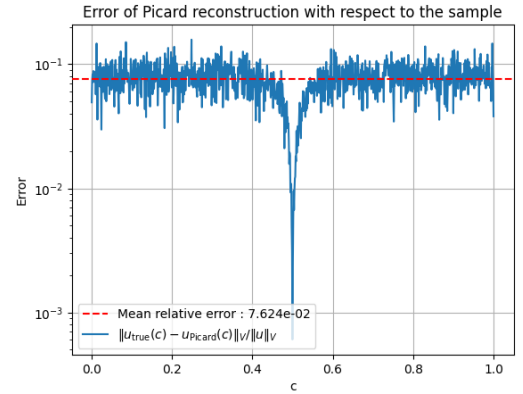
(a) Empirical Interpolation Method.



(b) Proper Orthogonal Decomposition.



(c) Gradient descent.



(d) Picard iteration method.

Figure 7: Error comparison for $\mu \in \mathbb{P}_{h,\text{test}}$, $m = N_{\text{modes}} = 2$, $\sigma = 10^{-1}$.

The Gradient Descent with Machine Learning and Picard iterations perform slightly better with a huge improvement for $\mu = 1/2$ (which is the root of the function $\alpha(\mu)$).

In Table 4, we showcase the average time taken to reconstruct 1 sample from $\mathbf{y}^\circ$.

Method	POD	Picard	EIM	Gradient Descent
Time (s/sample)	3.35×10^{-5}	3.83×10^{-4}	1.66×10^{-5}	3.27×10^{-2} .

Table 4: Computation time per sample of different methods for orthogonal modes.

We note that EIM is the fastest because the interpolating matrix is cached, and $n = \bar{n} = 1$. We think that POD or Picard iterations can be faster for $n \gg \bar{n}$. The real bottleneck seems to be the Gradient Descent which is slowed by the GPflow auto-differentiation of the GPR model.

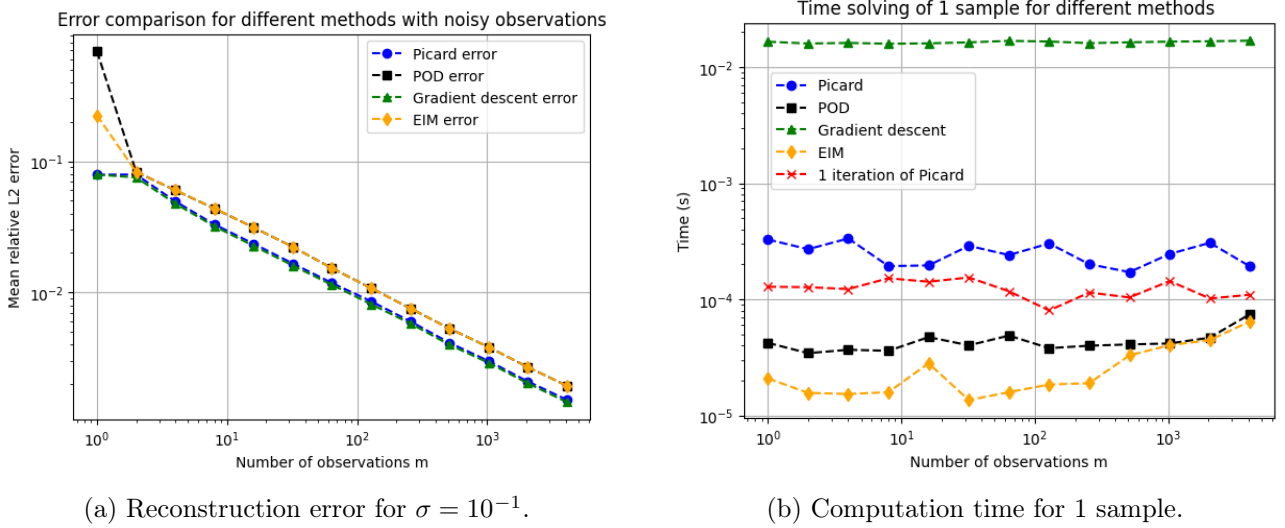


Figure 8: Performance comparison of different methods when the number of sensors $m = 2^i$, $i = 0, \dots, 12$ and $n = \bar{n} = 1$.

In Figure 8a we observe similar performances of all the methods, with slight advantages for methods using the GPR model. Figure 8b highlights faster performances for EIM up to a certain number of sensors m from which the EIM is not faster anymore.

The methods using the GPR seems have constant solving time for m if $m \leq 4096$, $n = \bar{n} = 1$.

We note that for $m = 1$ (which implies that only 1 mode is used for EIM), the GPR closure slightly outperforms the EIM and POD method while still being fast in the Picard case.

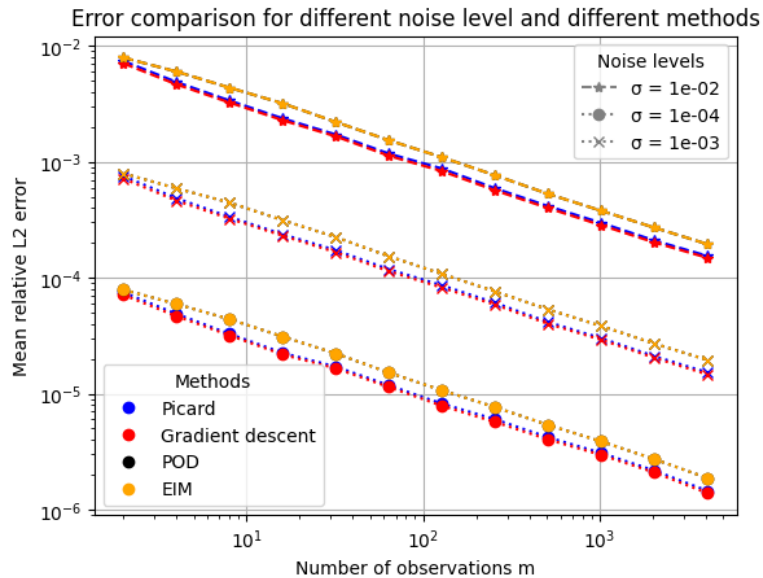


Figure 9: Relative error of reconstruction $\|\mathbf{u}_{n_d}(\mu) - \mathbf{u}^*(\mu)\|_X / \|\mathbf{u}_{n_d}\|_X$ depending on the noise level σ when $m = 2^i$, $i = 1, \dots, 12$ and $n = \bar{n} = 1$.

Figure 9 showcases the similar behavior regarding the relative noise level σ and the number of sensors m . Note we discarded $m = 1$ since it is a special case here.

3) Advection of a gaussian bump

a) Analysis of the POD learnability

We note that this test case is advection dominated, thus we expect the cumulative energy $\mathcal{E} = \sum_{i=1}^{\infty} \sigma_i^2$ to have slow decaying terms :

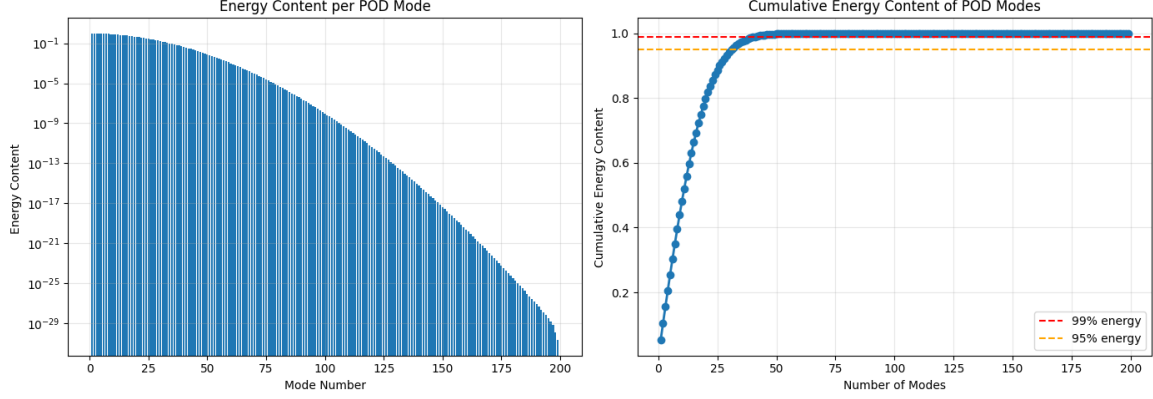


Figure 10: Energy histogram per modes (left). Cumulative energy $I = \sum_{i \leq n} \mathcal{E}_i / \sum_i \mathcal{E}_i$ (right).

We select value of the singular values σ_N is $\sigma_N \lesssim 10^{-16}$ which is numerical noise. We thus select $N = 199$.

	Number of Modes	MRE	RMSE	Relative L^∞ error
POD	$n = 199$	6.25×10^{-15}	3.08×10^{-14}	5.51×10^{-14}
EIM	$m = 199$	9.78×10^{-15}	2.88×10^{-12}	6.24×10^{-14}

Table 5: Error metrics of the EIM for different test cases in the $\|\cdot\|_X$ norm.

The following table compares in the higher part the error for different values of n with the same kernel initialisation.

The lower part compares different kernel that has been tried. Only Matern and RBF have been able to showcase significant learning results for the whole set of coefficients.

Kernel	$(n, \bar{n})$	MRE	RMSE	Relative L^∞ error
RBF + Cste**2 + White	(1,100)	1.06×10^{-00}	3.04×10^{-00}	8.13×10^{-01}
RBF + Cste**2 + White	(2,100)	8.03×10^{-08}	9.72×10^{-06}	4.83×10^{-05}
RBF + Cste**2 + White	(3,100)	9.42×10^{-08}	2.85×10^{-06}	5.69×10^{-05}
RBF + Cste**2 + White	(2,125)	2.62×10^{-06}	3.00×10^{-04}	1.04×10^{-03}
RBF + Cste**2 + White	(3,125)	2.82×10^{-06}	3.23×10^{-04}	1.13×10^{-03}
RBF + Cste**2 + White	(2,197)	5.98×10^{-05}	5.18×10^{-03}	2.98×10^{-02}
Matern($\nu=2.5$) + Cste**2 + White	(2,197)	1.24×10^{-04}	4.96×10^{-04}	1.68×10^{-02}
Matern($\nu=1.5$) + Cste**2 + White	(2,197)	1.56×10^{-03}	1.01×10^{-02}	3.02×10^{-02}

Table 6: Error metrics of the GPR output $h(\alpha(\mu))$ with respect to the target value $\beta(\mu)$ for different kernels and $\mu \in \mathbb{P}_{h,\text{test}}$.

The value coefficients α_i are expected to be regular in Appendix B.

Since $\text{RBF} = \text{Matern}(\nu = \infty)$ where ν controls the regularity of Matern functions, it is no surprise that the best results are obtained with RBF kernel or Matern with large ν .

Following Table 6, it seems that the α_i, β_i are indeed regular and the real dimension of the coefficients manifold is $n = 2$.

We also note that more modes to learn ($\bar{n}$ increases) translates into more to learn GPR learning error, which is no surprise and should be fixed with more numerical samples.

In the following graphs, we compare the error of the POD $\|\mathbf{u}_{n_d} - \mathbf{u}_{\text{POD}}\|_X = \|\boldsymbol{\eta}_{\text{POD}}\|_X$, the error of the ML constructed function $\|\mathbf{u}_{n_d} - \mathbf{u}_N\|_X$ and the learning error $\|\boldsymbol{\beta} - h(\boldsymbol{\alpha})\|_2$ for different trained GPR.

We use the Greedy quadratic error closure $\mathbf{u}_{n_d} - \mathbf{u}_{n,\text{POD}} = \boldsymbol{\eta}_{\text{GQuadMani}}$ proposed in [SP24; JAI+17] as a reference comparison for error modeling. We also verify Eq.(1.31) for $n = 2, \bar{n} = 123, 148$:

Approximation error depending on the chosen method with a GPR trained on $N=125$ modes

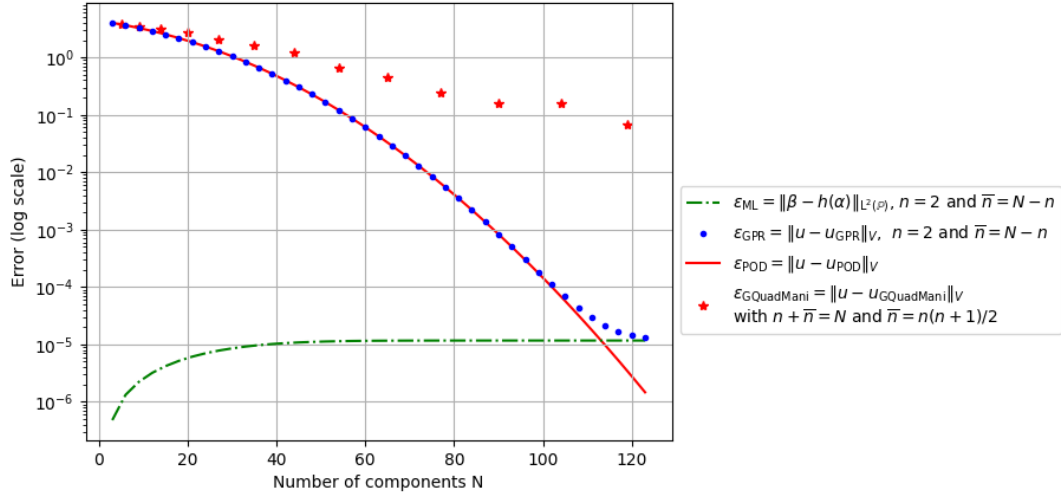


Figure 11: Error on reconstructing $\mathbf{u}_{n_d}$ with a GPR trained on $n = 2, \bar{n} = 125$.

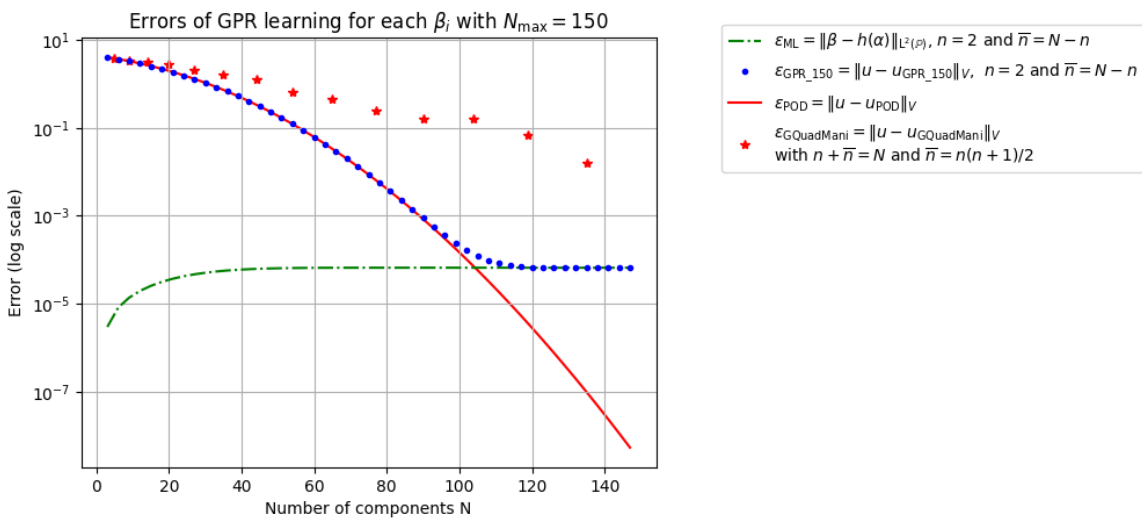


Figure 12: Error on reconstructing $\mathbf{u}_{n_d}$ with a GPR trained on $n = 2, \bar{n} = 150$.

The reconstruction error attains a plateau : this is due to the decay of the norms ($\|\beta_i\|_2 = \sigma_i$). Thus even if the relative error increases, the reconstruction error stays stable :

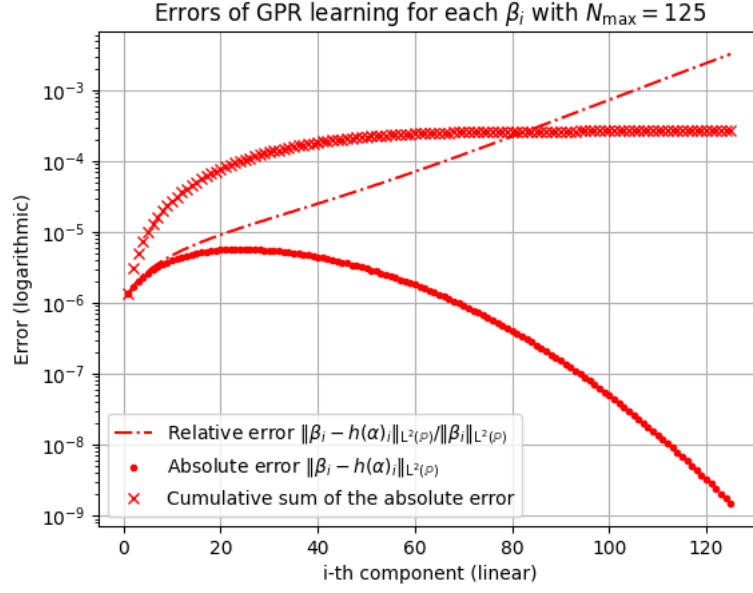


Figure 13: Different learning error metrics for a GPR trained with $n = 2$, $\bar{n} = 125$.

The relative error is linear in the loglog scale, translating the difficulty to learn the high frequencies of the POD coefficients with the low frequencies ones, see Figure 14 and Figure 15.

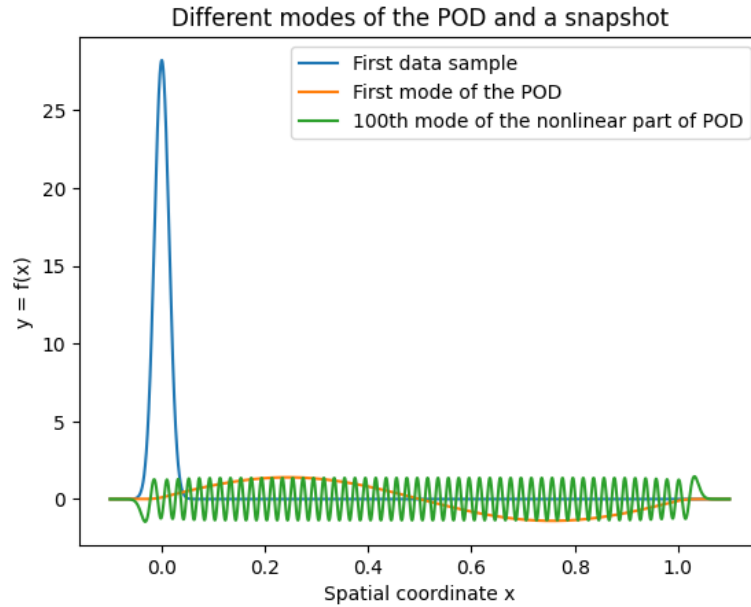


Figure 14: The first trajectory of the dataset, the first mode of the POD basis and the 100-th mode of the POD basis.

It is clear that the lower order modes are slowly oscillating while the higher order modes are fast oscillating corrections. Nyquist-Shannon sampling theorem states that we will not be able to recover highly oscillating coefficients by only sampling the first mode. Thus we expect $n > 1$.

Figure 15 shows that β_{100} is hard to describe only from α_1 . The trajectory $\beta_{100}(\mu)$ seems to be in a 2D space as the GPR learning showcased.

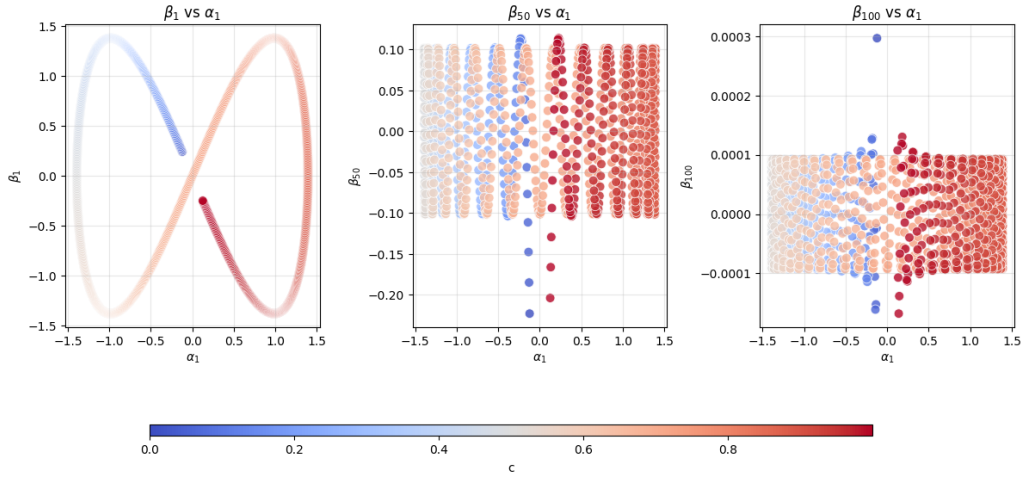


Figure 15: $\beta_i(\alpha_1)$ plotted as function of $\alpha_1(c)$ when the parameter $c \in \mathbb{P}$ varies and $i = 1, 50, 100$.

This can be verified plotting $\beta_{100} = \beta_{100}(\alpha_1, \alpha_2)$ in Figure 16 :

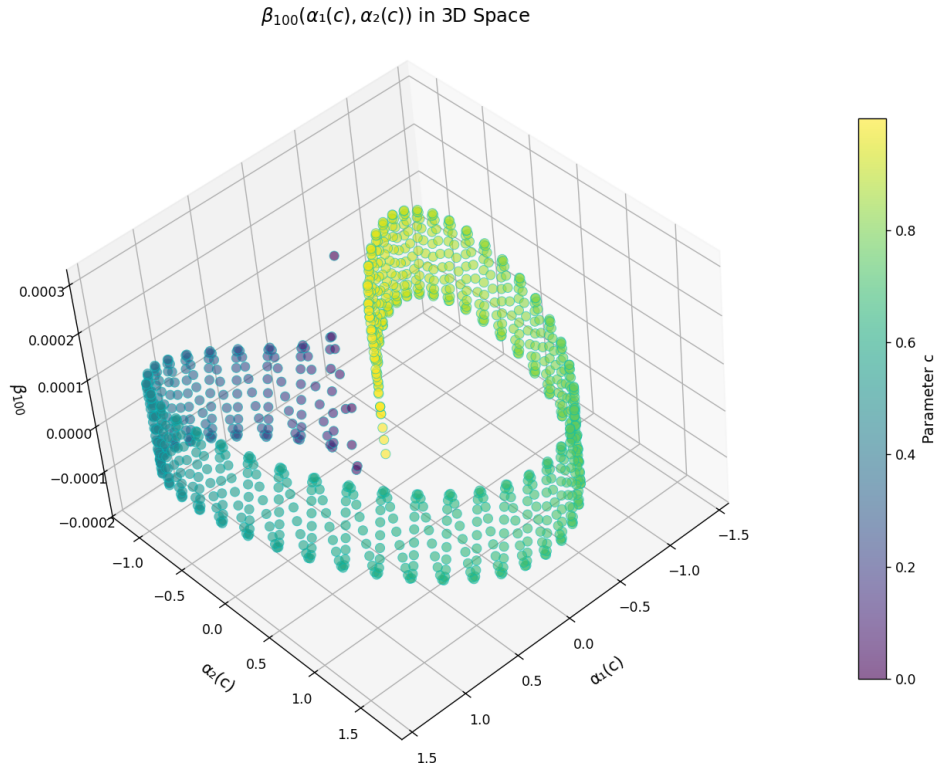


Figure 16: Coefficient $\beta_{100}(\alpha_1, \alpha_2)$ plotted for the different values of $\alpha_1(c), \alpha_2(c)$ when the parameter $c \in \mathbb{P}$ varies.

The results of learning of Table 6 are satisfying considering the measurements error level considered in industry range from $\sigma = 10^{-1}$ to $\sigma = 10^{-4}$ thus we need the ML model to allow representation up to a tolerance of 10^{-4} .

b) Online data assimilation

We consider $n = 2$ and $\bar{n} = 150$. The GPR Kernel is $\text{RBF}(\ell=0.0223) + 0.00188 \cdot 2 + \text{White}(\epsilon=1e-12)$.

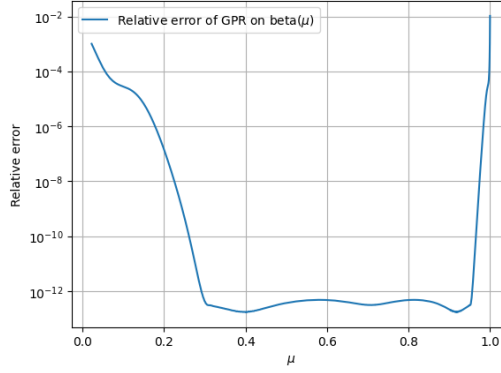
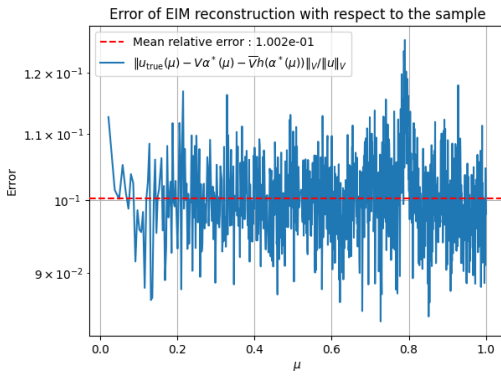


Figure 17: Relative error of the GPR over the samples $\mu \in \mathbb{P}_{h,\text{test}}$. Mean relative $\mathbf{L}^2$ -error : 5.98×10^{-05} .

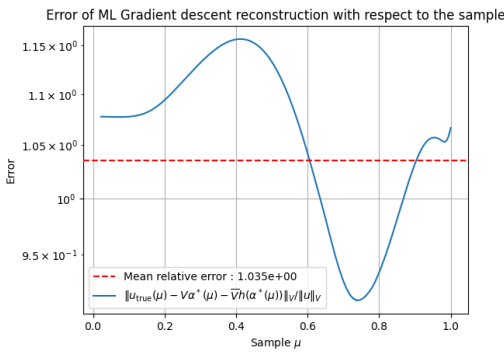
Figure 18 showcases the error on each samples for $\sigma = 10^{-1}$.



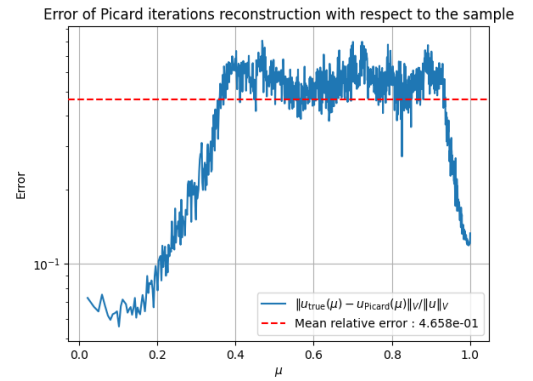
(a) Empirical Interpolation Method.



(b) Proper Orthogonal Decomposition.



(c) Gradient descent with Machine Learning.



(d) Picard iteration method.

Figure 18: Error comparison for $\mu \in \mathbb{P}_{h,\text{test}}$, $m = N_{\text{modes}} = 2$, $\sigma = 10^{-1}$. Advection case.

For the error, a plain Least Square on the POD coefficients seems to perform the best, while Picard or Gradient Descent don't seem to perform better. This may be explain in the further discussion. Note that Picard iterations converge usually after two or three iterations at most.

Method	POD	Picard	EIM	Gradient Descent
Time (s/sample)	6.67×10^{-4}	1.63×10^{-3}	5.56×10^{-4}	1.31×10^0 .
MRE	5.69×10^{-2}	4.27×10^{-1}	1.00×10^{-1}	1.035×10^0 .

Table 7: Comparison of different methods for advection.

The time for POD and EIM are alike in term of computation time while Picard corresponds to a few EIM computation. This is no surprise since Picard is an iteration of Least Square + ML evaluation. The bottleneck is on the Gradient Descent : the GPflow implementation applies autodifferentiation on 150 independant GPR with a cost of evaluation of ~ 1 seconds per sample. This is not tractable.

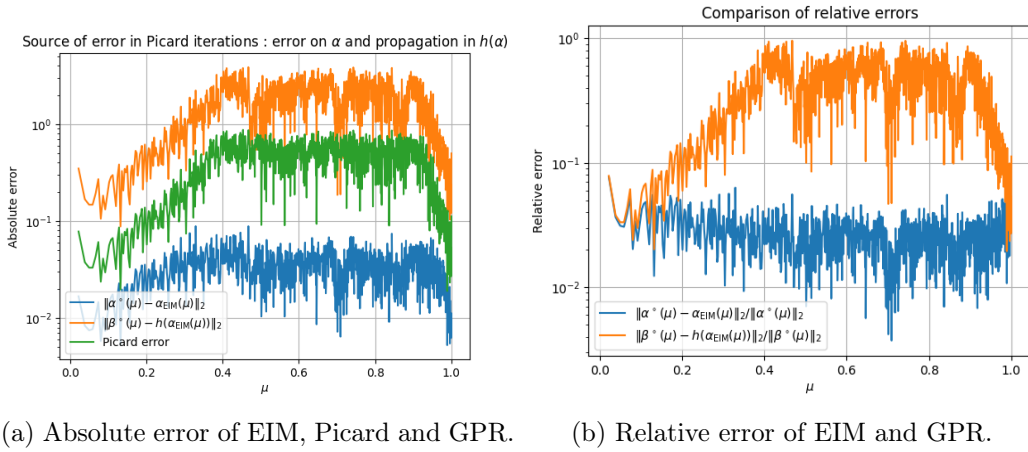
Figure 19: L^2 -error for each samples of the GPR, EIM and Picard data assimilation method.

Figure 19 highlights the propagation of the error on α° from α_{EIM} through the GPR model h .

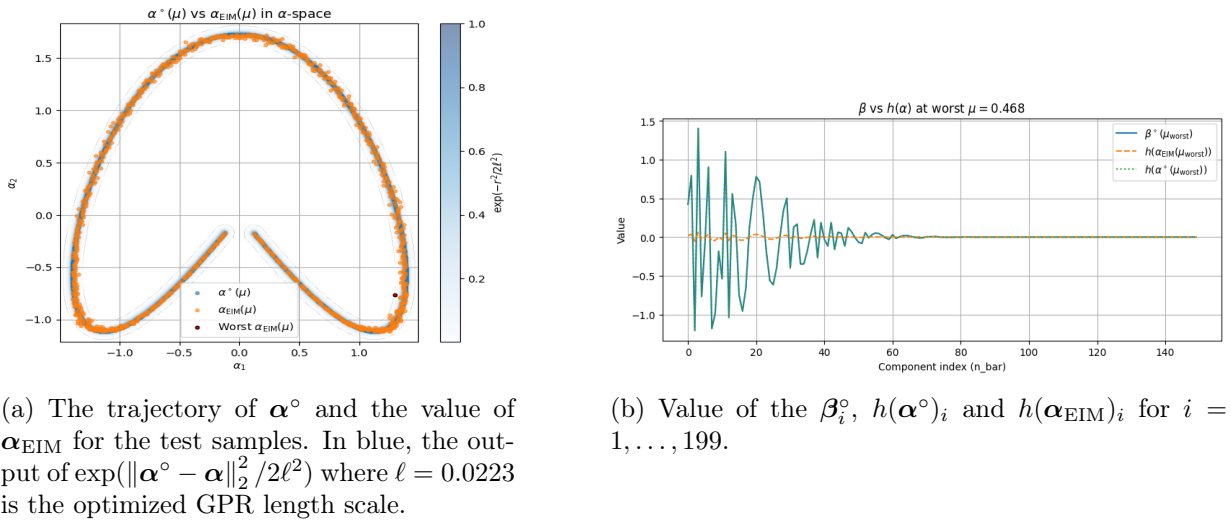


Figure 20: Graphic analysis of the GPR sensitivity to noise.

Figure 20 shows that a GPR trained on non noisy α with high confidence may return a GPR with length scale ℓ too small compared to the scale of the noise. It results as a GPR returning almost null values for noisy observations!

Some conclusions can be drawn out of these results :

1. The GPR **must** be trained on a noisy dataset or with length scale adapted to the noise.
2. Considering more coefficients β_i to learn does not necessarily mean a better approximation.
3. The method of optimization must be **gradient-free** or parallel evaluations of the GPR should be used.
4. For $m = n_{\text{EIM}}$, the Picard fixed-point and Gradient Descent procedures did **NOT** improve much the accuracy with respect to the initialization, thus no need for online optimization after evaluation the first coefficients.

Answering the first conclusion will not be an issue and only needs a retraining of the GPR.

Answering the third conclusion has already being done with Picard iteration. We can also consider setting $m = N_{\text{EIM}}$ and consider, having evaluated $\alpha \in \mathbb{R}^n$,

$$\tilde{\mathbf{u}} = \mathbf{u}_{\text{ref}} + \mathbf{V}\alpha + \overline{\mathbf{V}}h(\alpha)$$

with no optimization process.

It is also possible to consider new maps $y \mapsto (\gamma_i)_{1 \leq i \leq n}, (\hat{u}_i)_{1 \leq i \leq n} \mapsto (\gamma_i)_{i > n}, (\hat{u}_i)_{i > n}$.

For the fourth conclusion, it is actually easier to only evaluate the mapping once instead of optimizing over it.

For the second conclusion, it is important to take into consideration that the number of POD modes N depends on the sampling i.e. $N = N_s(N_{\text{train}})$. Thus depending on the number of samples considered, the number of physical reduced basis modes could be reduced and the efforts could be concentrated into learning only meaningful coefficients.

It is clear that the lower order modes are slowly oscillating while the higher order modes are fast oscillating corrections. Nyquist-Shannon sampling theorem states that we will not be able to recover highly oscillating coefficients by only sampling the first mode. Thus we expect $n > 1$.

Conclusion

In this work, the encoding of reduced basis methods using Gaussian Process Regressors for data assimilation was studied.

It was demonstrated the capacity of Machine Learning to discover the latent space dimension of the reduced basis coefficients and the possibility to encode a full reduced basis into only a few modes, around the number of parameters.

Several data assimilation methods have been proposed to reconstruct the full field from partial observations. The Empirical Interpolation Method has been identified as a good initialization point when needed, and has also served as benchmark to compare the results.

Gradient methods were shown to slow down the data assimilation process for slight to no improvements while fixed-points methods or direct evaluation of the encoder could prove to beat the EIM.

It was also highlighted that if the latent dimension of the reduced basis was small enough with respect to the number of basis functions and the number of observations, the evaluation time are equivalent between EIM and the fixed-point method.

Last it is needed to construct the encoder with scales comparable to the relative noise scale and not enough to just learn the exact mapping.

Sampling considerations were done and need to be further studied in order to improve the selection of the reduced basis modes, thus the final reconstruction precision.

Some perspectives consist of focusing on the different mappings learning from the observation to the reconstructed fields by learning e.g. the map from EIM coefficients to POD, or EIM to the difference between EIM and POD coefficients and the last POD coefficients.

These could be used to reconstruct the field without minimization on α but just reconstructing the full field directly from the first coefficients and the evaluation of the ML model.

Part 3

Appendix

A Error metrics

1) Parameter norm

Let $P \geq 1$ be the parameter dimension. Let $\mathbb{P} \subset \mathbb{R}^P$.

Let N_s be the number of samples used to represent the parameter space $\mathbb{P}$.

Let $n \geq 1$ be the dimension of an output of interest. Let $g : \boldsymbol{\mu} \in \mathbb{P} \mapsto g(\boldsymbol{\mu}) \in \mathbb{R}^n$ be the output function. For $\boldsymbol{\mu} \in \mathbb{P}$, the 2-norm of $g(\boldsymbol{\mu}) \in \mathbb{R}^n$ is :

$$\|g(\boldsymbol{\mu})\|_{2,\mathbb{R}^n}^2 = \sum_{i=1}^n |(g(\boldsymbol{\mu}))_i|^2.$$

Following [QSS06, Section 9.1], let $\boldsymbol{\mu}_1, \dots, \boldsymbol{\mu}_{N_s}$ be quadrature points in $\mathbb{P}$ and let $w_1, \dots, w_{N_s} > 0$ be quadrature weights.

The discrete $\mathbf{L}^2(\mathbb{P})$ norm is :

$$\|g(\boldsymbol{\mu})\|_{\mathbf{L}^2(\mathbb{P}),h}^2 = \sum_{i=1}^{N_s} w_i \|g(\boldsymbol{\mu}_i)\|_{2,\mathbb{R}^n}^2. \quad (3.1)$$

When the quadrature is uniform, this reduce to $w_i = \frac{1}{N_s}$:

$$\|g(\boldsymbol{\mu})\|_{\mathbf{L}^2(\mathbb{P}),h}^2 = \frac{1}{N_s} \sum_{i=1}^{N_s} \|g(\boldsymbol{\mu}_i)\|_{2,\mathbb{R}^n}^2.$$

For $\boldsymbol{\mu} \in \mathbb{P}$, the supremum norm of $g(\boldsymbol{\mu}) \in \mathbb{R}^n$ is :

$$\|g(\boldsymbol{\mu})\|_{\infty,\mathbb{R}^n} = \max_{1 \leq i \leq n} |g(\boldsymbol{\mu})_i|.$$

The supremum norm over the parameters is defined as :

$$\|g(\boldsymbol{\mu})\|_{\mathbf{L}^\infty(\mathbb{P})} = \max_{\boldsymbol{\mu}_i \in \mathbb{P}_h} \|g(\boldsymbol{\mu}_i)\|_{\infty,\mathbb{R}^n}. \quad (3.2)$$

2) Timespace norm

a) Spatial norm

Let $d \geq 1$ be the space dimension. Let $\Omega \subset \mathbb{R}^d$ be an open domain. Let $n_d \geq 1$ be the number of nodes discretizing Ω . Let $V \subset \mathbf{L}^2(\Omega, \mathbb{R}^{n_d})$ be an Hilbert.

Let $\mathbf{M}_\delta$ be a mass matrix that can be constructed either by quadrature points over the spatial domain Ω or the basis functions of the Finite Elements space.

If $\mathbf{u}_{n_d} \in V_h$, the 2-norm is defined as :

$$\|\mathbf{u}_{n_d}\|_{2,\mathbb{R}^{n_d}}^2 = \sum_{i=1}^{n_d} |(\mathbf{u}_{n_d})_i|^2.$$

If $\mathbf{u}_{n_d} \in V_h$, the energy norm or V -norm is defined as :

$$\|\mathbf{u}_{n_d}\|_V^2 = \left\| \mathbf{M}_\delta^{1/2} \mathbf{u}_{n_d} \right\|_{2,\mathbb{R}^{n_d}}^2 = \sum_{i=1}^{n_d} \left| (\mathbf{M}_\delta^{1/2} \mathbf{u}_{n_d})_i \right|^2. \quad (3.3)$$

The supremum norm of $\mathbf{u}_{n_d} \in \mathbb{R}^{n_d}$ is:

$$\|\mathbf{u}_{n_d}\|_{\infty,\mathbb{R}^{n_d}} = \max_{1 \leq i \leq n_d} |(\mathbf{u}_{n_d})_i|. \quad (3.4)$$

b) Timespace norm

Let $d \geq 1$ be the space dimension. Let $\Omega \subset \mathbb{R}^d$ be an open domain. Let $n_d \geq 1$ be the number of nodes discretizing Ω .

Let $T > 0$, let $I = [0; T]$ be a time interval. Let $V \subset \mathbf{L}^2(\Omega, \mathbb{R}^{n_d})$ be an Hilbert.

Let $\mathbf{M}_\delta$ be a mass matrix that can be constructed either by quadrature points over the spatial domain Ω or by the $\mathbf{L}^2$ -innerproduct of the Finite Elements basis functions.

If $\mathbf{u}_{n_d} \in \mathbf{L}^2(I, V)$, the energy norm or V -norm is defined as :

$$\|\mathbf{u}_{n_d}\|_V^2 = \int_I \left\| \mathbf{M}_\delta^{1/2} \mathbf{u}_{n_d}(t) \right\|_{2, \mathbb{R}^{n_d}}^2 dt. \quad (3.5)$$

If I is uniformly discretized in $N_{\text{timesteps}} \geq 1$, $0 = t_0 < t_1 < \dots < t_{N_{\text{timesteps}}} = T$, this reduce to :

$$\|\mathbf{u}_{n_d}\|_V^2 = \frac{1}{N_{\text{timesteps}} + 1} \sum_{i=0}^{N_{\text{timesteps}}} \left\| \mathbf{M}_\delta^{1/2} \mathbf{u}(t_i) \right\|_{2, \mathbb{R}^{n_d}}^2. \quad (3.6)$$

The supremum norm of $\mathbf{u}_{n_d} \in \mathbf{L}^2(I, \mathbb{R}^{n_d})$ is :

$$\|\mathbf{u}_{n_d}\|_{\infty, \mathbf{L}^2(I, \mathbb{R}^{n_d})} = \max_{1 \leq i \leq N_{\text{timesteps}}} \|\mathbf{u}_{n_d}(t_i)\|_{\infty, \mathbb{R}^{n_d}}. \quad (3.7)$$

3) Error metrics

Let $(E, \|\cdot\|_E)$ be a normed space. Let $y \in E$ be a target value. Let $v \in E$ be an approximation of y .

The error in norm $\|\cdot\|_E$ between y and v is :

$$\text{Err}_E(v, y) := \|y - v\|_E. \quad (3.8)$$

The relative error in norm $\|\cdot\|_E$ between y and v is :

$$\text{RelErr}_E(v, y) := \begin{cases} \|y - v\|_E / \|y\|_E & \text{if } y \neq 0, \\ \|v\|_E & \text{if } y = 0. \end{cases} \quad (3.9)$$

Let $(y_i)_{1 \leq i \leq N_s} \in (R^n)^{N_s}$ is a sample of N_s target in R^n and $(v_i)_{1 \leq i \leq N_s}$ is a sample of N_s approximations, let $w_i > 0$ be quadrature weights.

The Mean Relative Error (RME) is :

$$\text{MRE}_E(\mathbf{y}, \mathbf{v}) := \sum_{i=1}^{N_s} w_i \text{RelErr}_E(v_i, y_i) / \sum_i w_i. \quad (3.10)$$

In the uniform sampling case, this reduce to :

$$\text{MRE}_E(\mathbf{y}, \mathbf{v}) := \frac{1}{N_s} \sum_{i=1}^{N_s} \text{RelErr}_E(v_i, y_i). \quad (3.11)$$

The Root Mean Square Error (RMSE) is their average $\mathbf{L}^2$ -norm over the sampling :

$$\text{RMSE}_{2, \mathbb{R}^n}(\mathbf{y}, \mathbf{v}) := \sqrt{\sum_{i=1}^{N_s} w_i \|y_i - v_i\|_{2, \mathbb{R}^n}^2 / \sum_i w_i}. \quad (3.12)$$

In the uniform sampling case, this reduce to :

$$\text{RMSE}_{2, \mathbb{R}^n}(\mathbf{y}, \mathbf{v}) := \sqrt{\frac{1}{N_s} \sum_{i=1}^{N_s} \|y_i - v_i\|_{2, \mathbb{R}^n}^2}.$$

B Test cases construction

1) Orthogonal modes

Let $I = [0; \pi]$, $\mathbb{P} = [0; 1]$, $\lambda > 0$.

We consider the function parameterized by $\mu \in \mathbb{P}$:

$$u(x ; \mu) = \alpha_1(\mu)\Phi_1(x) + \lambda\alpha_2(\mu)\Phi_2(x). \quad (3.13)$$

Where we construct Φ_1, Φ_2 and α_1, α_2 such that :

$$\begin{aligned} \langle \Phi_i | \Phi_j \rangle_{L^2(I)} &= \delta_{ij}, \\ \langle \alpha_i | \alpha_j \rangle_{L^2(\mathbb{P})} &= \delta_{ij}, \\ \int_{\mathbb{P}} \alpha_1 &= \int_{\mathbb{P}} \alpha_2 = 0. \end{aligned} \quad (3.14)$$

We want to force $\Phi_1 = \sin / \|\sin\|_{L^2(I)}$ and $\Phi_2 = \cos / \|\cos\|_{L^2(I)}$.

To force $\Phi_2 = \sqrt{\frac{2}{\pi}} \cos$ to be the 2nd POD mode, we set $\lambda = 0.1$ such that $\cos$ will represent exactly 1/100 of the total energy.

Note that

$$\operatorname{argmax}_{\varphi \in \operatorname{Span}(\Phi_1, \Phi_2)} \|\varphi\|_{L^\infty(\Omega)} = \{\Phi_1, \Phi_2\}$$

and

$$\operatorname{argmax}_{x \in [0; \pi]} |\Phi_1(x)| = \left\{ \frac{\pi}{2} \right\}, \quad \operatorname{argmax}_{x \in [0; \pi]} |\Phi_2| = \{0, \pi\}.$$

Thus expect the EIM to select $x_1 \in \{0, \frac{\pi}{2}, \pi\}$ as first sensor locations, then choose $x_2 = \frac{\pi}{2}$ if $x_1 \neq \frac{\pi}{2}$ or $x_2 \in \{0, \pi\}$ if $x_1 = \frac{\pi}{2}$.

2) Accelerating advection of a gaussian bump

Let $I = [-0.1; 1.1] = [a; b]$, $\mathbb{P} = [0; 1]$ ($P = 1$). Let $s > 0$.

We consider an accelerating advection of a gaussian bump :

$$u(x ; \mu) = \frac{1}{(2\pi s^2)^{1/2}} \exp\left(-\frac{(x - \mu^2)^2}{2s^2}\right). \quad (3.15)$$

We choose $s^2 = 0.01$ such that $u(x = -0.1 ; \mu = 0) = u(x = 1.1 ; \mu = 1) \simeq 0$.

We expand the square to separate the coefficient from the spatial parameter :

$$u(x ; \mu) = \frac{1}{(2\pi s^2)^{1/2}} \exp\left(\frac{-x^2}{2s^2}\right) \exp\left(\frac{x\mu}{s^2}\right) \exp\left(\frac{-\mu^4}{2s^2}\right).$$

The middle part can be written with the power serie of exponential :

$$u(x ; \mu) = \sum_{n=0}^{\infty} e^{\frac{-\mu^4}{2s^2}} \frac{\mu^{2n} x^n}{s^{2n} n!} e^{\frac{-x^2}{2s^2}} = \sum_{n=0}^{\infty} \alpha_n(\mu) q_n(x).$$

The family of functions $q_n(x) := x^n e^{\frac{-x^2}{2s^2}}$ is linearly independent. Thus $(q_n)_n$ is an infinite dimensional family of $L^2(I)$. Hence the function u can not be represented by a finite dimensional basis, otherwise there would exist a finite dimensional basis of $L^2(I)$ that contains an infinite dimensional basis of $L^2(I)$.

We conclude that the numerical POD basis can have infinite dimension up to numerical precision.

3) Burgers equation

This part collects an analysis of Burgers equation made during the internship that we did not have time to try regarding the difficulties encountered on the advection case.

a) The continuous problem

Let $\Omega =]0; 1]^2$ be a space domain. Let $I = [0; 1]$ be a time interval. Let $\nu \geq 0$ be a viscosity parameter. Let $\delta_1, \delta_2 > 0$ be small enough to respect Eq.(3.35), $\mathbb{P} \subset [1 - \delta_1; 1 + \delta_1] \times [0.1 - \delta_2; 0.1 + \delta_2]$. Let $\boldsymbol{\mu} = (a, s) \in \mathbb{P}$.

We note the initial condition u_0 as :

$$u_0(\mathbf{x} ; \boldsymbol{\mu}) = \frac{1}{2\pi\sigma^2} \exp \left(-\frac{(x_1 - a)^2 + (x_2 - a)^2}{2s^2} \right).$$

We consider the solution of the Burgers equation :

$$\text{Find } u(\cdot, \cdot ; \boldsymbol{\mu}) \text{ s.t. } \begin{cases} \frac{\partial u}{\partial t} + \operatorname{div} \left(\frac{u^2}{2} \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right) - \nu \Delta u & = & 0 & \text{in } \Omega, \\ u(\mathbf{x}, 0) & = & u_0(\mathbf{x} ; \boldsymbol{\mu}) & \text{in } \Omega, \\ u(\mathbf{x}, t) & = & 0 & \text{on } \partial\Omega. \end{cases} \quad (3.16)$$

If $\nu = 0$, the expected solution is :

$$u(x, t) = u_0(x + u_0(x, y)t).$$

Since $u_0 > 0$ the solution is a advected gaussian bump along the line $\mathbb{R} \begin{bmatrix} 1 \\ 1 \end{bmatrix}$ from $(x_0, y_0) = (a, a)$.

If $\nu > 0$, then the solution follow the same advection trajectory with dispersion of the energy.

b) Discrete equation

We rewrite this in variational form.

For $u \in L^2(\mathbb{R}_+^*; H_0^1(\Omega))$:

$$\begin{aligned} & \frac{\partial u}{\partial t} + \operatorname{div} \left(\frac{u^2}{2} \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right) - \nu \Delta u & = & 0 \\ \iff \forall v \in H_0^1(\Omega), & \left\langle \frac{\partial u}{\partial t} \middle| v \right\rangle_{L^2(\Omega)} + \left\langle \operatorname{div} \left(\frac{u^2}{2} \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right) \middle| v \right\rangle_{L^2(\Omega)} - \nu \langle \Delta u | v \rangle_{L^2(\Omega)} & = & 0 \\ \iff \forall v \in H_0^1(\Omega), & \left\langle \frac{\partial u}{\partial t} \middle| v \right\rangle_{L^2(\Omega)} - \left\langle \frac{u^2}{2} \begin{pmatrix} 1 \\ 1 \end{pmatrix} \middle| \nabla v \right\rangle_{L^2(\Omega)} + \nu \langle \nabla u | \nabla v \rangle_{L^2(\Omega)} & = & 0. \end{aligned}$$

We discretize in time with finite difference using implicit diffusion and semi-implicit advection schemes:

$$\begin{aligned} & \forall v \in H_0^1(\Omega), \left\langle \frac{u^{n+1} - u^n}{dt} \middle| v \right\rangle_{L^2(\Omega)} - \left\langle \frac{u^n u^{n+1}}{2} \begin{pmatrix} 1 \\ 1 \end{pmatrix} \middle| \nabla v \right\rangle_{L^2(\Omega)} + \nu \langle \nabla u^{n+1} | \nabla v \rangle_{L^2(\Omega)} = 0 \\ \iff \forall v \in H_0^1(\Omega), & \left\langle \frac{u^{n+1}}{dt} \middle| v \right\rangle_{L^2(\Omega)} + \nu \langle \nabla u^{n+1} | \nabla v \rangle_{L^2(\Omega)} = \left\langle \frac{u^n}{dt} \middle| v \right\rangle_{L^2(\Omega)} + \left\langle \frac{u^n u^{n+1}}{2} \begin{pmatrix} 1 \\ 1 \end{pmatrix} \middle| \nabla v \right\rangle_{L^2(\Omega)} \end{aligned}$$

We then use the finite element approximation for a FEM basis $(\Phi_i)_i$ such that for all n , $u^n = \sum_i u_i^n \Phi_i$ and let $v = \Phi_j$:

$$\sum_i \left(\frac{u_i^{n+1}}{dt} \langle \Phi_i | \Phi_j \rangle_{L^2(\Omega)} + u_i^{n+1} \nu \langle \nabla \Phi_i | \nabla \Phi_j \rangle_{L^2(\Omega)} \right) \quad (3.17)$$

$$= \sum_i \frac{u_i^n}{dt} \langle \Phi_i | \Phi_j \rangle_{L^2(\Omega)} + \left\langle \frac{(\sum_i u_i^n \Phi_i)(\sum_\ell u_\ell^{n+1} \Phi_\ell)}{2} \left| \frac{\partial \Phi_j}{\partial x} + \frac{\partial \Phi_j}{\partial y} \right\rangle_{L^2(\Omega)} \right. \quad (3.18)$$

This being true for all j we expand the last term as :

$$\frac{1}{2} \sum_{i\ell j} u_i^{n+1} u_\ell^n \left\langle \phi_i \phi_\ell \left| \frac{\partial \Phi_j}{\partial x} + \frac{\partial \Phi_j}{\partial y} \right\rangle_{L^2(\Omega)} \right. \quad (3.19)$$

We define :

$$C_{i\ell j} = \frac{1}{2} \left\langle \phi_i \phi_\ell \left| \frac{\partial \Phi_j}{\partial x} + \frac{\partial \Phi_j}{\partial y} \right\rangle_{L^2(\Omega)} \right. \quad (3.20)$$

$$N(U^n)_{\ell j} = \sum_i C_{i\ell j} u_i^n. \quad (3.21)$$

Thus the algebraic formulation :

$$\left(\frac{1}{dt} M_\delta + \nu K_\delta + N(U^n) \right) U^{n+1} = \frac{1}{dt} M_\delta U^n. \quad (3.22)$$

$$(3.23)$$

We replace now u^n by its approximation $\bar{U} + \Phi \alpha^n$:

$$\left(\frac{1}{dt} M_\delta + \nu K_\delta + N(\bar{U} + \Phi \alpha^n) \right) (\bar{U} + \Phi \alpha^{n+1}) = \frac{1}{dt} M_\delta (\bar{U} + \Phi \alpha^n) \quad (3.24)$$

$$\iff \left(\frac{1}{dt} M_\delta + \nu K_\delta + N(\bar{U} + \Phi \alpha^n) \right) \Phi \alpha^{n+1} = \frac{1}{dt} M_\delta \Phi \alpha^n - (\nu K_\delta - N(\bar{U} + \Phi \alpha^n)) \bar{U}. \quad (3.25)$$

$$\iff \left(\frac{1}{dt} \Phi^T M_\delta + \nu K_\delta + N(\bar{U} + \Phi \alpha^n) \right) \Phi \alpha^{n+1} = \frac{1}{dt} M_\delta \Phi \alpha^n - (\nu K_\delta - N(\bar{U} + \Phi \alpha^n)) \bar{U}. \quad (3.26)$$

We use Eq.(3.26) as the algebraic formulation of the PDE solver.

This formulation is stable under a CFL condition :

$$\forall v \in H_0^1(\Omega), \left\langle \frac{u^{n+1}}{dt} \middle| v \right\rangle_{L^2(\Omega)} + \nu \langle \nabla u^{n+1} | \nabla v \rangle_{L^2(\Omega)} = \left\langle \frac{u^n}{dt} \middle| v \right\rangle_{L^2(\Omega)} + \left\langle \frac{u^n u^{n+1}}{2} \begin{pmatrix} 1 \\ 1 \end{pmatrix} \middle| \nabla v \right\rangle_{L^2(\Omega)} \quad (3.27)$$

$$v = u^{n+1} \implies \left\langle \frac{u^{n+1}}{dt} \middle| u^{n+1} \right\rangle_{L^2(\Omega)} + \nu \langle \nabla u^{n+1} | \nabla u^{n+1} \rangle_{L^2(\Omega)} = \left\langle \frac{u^n}{dt} \middle| u^{n+1} \right\rangle_{L^2(\Omega)} + \left\langle \frac{u^n u^{n+1}}{2} \begin{pmatrix} 1 \\ 1 \end{pmatrix} \middle| \nabla u^{n+1} \right\rangle_{L^2(\Omega)} \quad (3.28)$$

$$\iff \left\langle \frac{u^{n+1} - u^n}{dt} \middle| u^{n+1} \right\rangle_{L^2(\Omega)} + \nu \|\nabla u^{n+1}\|_{L^2(\Omega)}^2 = \left\langle \frac{u^n u^{n+1}}{2} \begin{pmatrix} 1 \\ 1 \end{pmatrix} \middle| \nabla u^{n+1} \right\rangle_{L^2(\Omega)}. \quad (3.29)$$

$$(3.30)$$

We use the following equality :

$$\left\langle \frac{u^{n+1} - u^n}{dt} \middle| u^{n+1} \right\rangle_{\mathbf{L}^2(\Omega)} = \frac{1}{2dt} \left(\|u^{n+1}\|_{\mathbf{L}^2(\Omega)}^2 - \|u^n\|_{\mathbf{L}^2(\Omega)}^2 + \|u^{n+1} - u^n\|_{\mathbf{L}^2(\Omega)}^2 \right).$$

It gives the following estimate :

$$\frac{1}{2dt} \left(\|u^{n+1}\|_{\mathbf{L}^2(\Omega)}^2 - \|u^n\|_{\mathbf{L}^2(\Omega)}^2 + \|u^{n+1} - u^n\|_{\mathbf{L}^2(\Omega)}^2 \right) + \nu \|\nabla u^{n+1}\|_{\mathbf{L}^2(\Omega)}^2 = \left\langle \frac{u^n u^{n+1}}{2} \begin{pmatrix} 1 \\ 1 \end{pmatrix} \middle| \nabla u^{n+1} \right\rangle_{\mathbf{L}^2(\Omega)}. \quad (3.31)$$

Then we focus on the nonlinear term of the right-hand side. Suppose $u^n \in L^\infty(\Omega)$:

$$\begin{aligned} \left\langle \frac{u^n u^{n+1}}{2} \begin{pmatrix} 1 \\ 1 \end{pmatrix} \middle| \nabla u^{n+1} \right\rangle_{\mathbf{L}^2(\Omega)} &= \left\langle \frac{u^n u^{n+1}}{2} \begin{pmatrix} 1 \\ 1 \end{pmatrix}^T \cdot \nabla u^{n+1} \right\rangle_{\mathbf{L}^2(\Omega)} \\ &\leq \frac{1}{2} \|u^n\|_{L^\infty(\Omega)} \|u^{n+1}\|_{\mathbf{L}^2(\Omega)} \underbrace{\left\| \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right\|_{\mathbb{R}^2, \infty}}_{=1} \|\nabla u^{n+1}\|_{\mathbf{L}^2(\Omega)}, \quad \text{Cauchy-Schwarz [BRE11]} \\ &\leq \left(\frac{1}{2\sqrt{\nu}} \|u^n\|_{L^\infty(\Omega)} \|u^{n+1}\|_{\mathbf{L}^2(\Omega)} \right) \left(\sqrt{\nu} \|\nabla u^{n+1}\|_{\mathbf{L}^2(\Omega)} \right) \\ &\leq \frac{1}{4\nu} \|u^n\|_{L^\infty(\Omega)}^2 \|u^{n+1}\|_{\mathbf{L}^2(\Omega)}^2 + \frac{1}{2} \nu \|\nabla u^{n+1}\|_{\mathbf{L}^2(\Omega)}^2, \quad \text{Young's inequality [BRE11]}. \end{aligned}$$

Thus we have the following inequality by inserting the last line to Eq.(3.31) :

$$\frac{1}{2dt} \left(\|u^{n+1}\|_{\mathbf{L}^2(\Omega)}^2 - \|u^n\|_{\mathbf{L}^2(\Omega)}^2 + \|u^{n+1} - u^n\|_{\mathbf{L}^2(\Omega)}^2 \right) + \nu \|\nabla u^{n+1}\|_{\mathbf{L}^2(\Omega)}^2 \quad (3.32)$$

$$\leq \frac{1}{4\nu} \|u^n\|_{L^\infty(\Omega)}^2 \|u^{n+1}\|_{\mathbf{L}^2(\Omega)}^2 + \frac{1}{2} \nu \|\nabla u^{n+1}\|_{\mathbf{L}^2(\Omega)}^2. \quad (3.33)$$

We rearrange the result :

$$\|u^{n+1}\|_{\mathbf{L}^2(\Omega)}^2 \left(1 - \frac{dt}{2\nu} \|u^n\|_{L^\infty(\Omega)}^2 \right) + \underbrace{\frac{\nu}{2} \|\nabla u^{n+1}\|_{\mathbf{L}^2(\Omega)}^2}_{0 \leq} \leq \|u^n\|_{\mathbf{L}^2(\Omega)}^2 - \underbrace{\|u^{n+1} - u^n\|_{\mathbf{L}^2(\Omega)}^2}_{\leq 0}.$$

Thus we have the following estimate linking $\|u^n\|_{\mathbf{L}^2(\Omega)}$ and $\|u^{n+1}\|_{\mathbf{L}^2(\Omega)}$:

$$\|u^{n+1}\|_{\mathbf{L}^2(\Omega)}^2 \left(1 - \frac{dt}{2\nu} \|u^n\|_{L^\infty(\Omega)}^2 \right) \leq \|u^n\|_{\mathbf{L}^2(\Omega)}^2. \quad (3.34)$$

This estimate prove the scheme is stable under the condition :

$$\begin{aligned} 0 &\leq 1 - \frac{dt}{2\nu} \|u^n\|_{L^\infty(\Omega)}^2 \\ \iff \frac{dt}{2\nu} \|u^n\|_{L^\infty(\Omega)}^2 &\leq 1 \\ \Rightarrow \frac{dt(dx)^2}{2\nu} \|u^n\|_{\mathbf{L}^2(\Omega)}^2 &\leq 1, \quad \text{inverse discrete inequality [AG21]}. \end{aligned}$$

Note that if

$$\frac{dt}{2\nu} \|u^n\|_{L^\infty(\Omega)}^2 \leq 1,$$

then

$$\|u^{n+1}\|_{\mathbf{L}^2(\Omega)} \leq \|u^n\|_{\mathbf{L}^2(\Omega)}.$$

Thus the CFL condition in essential supremum norm $\|\cdot\|_{L^\infty}$ is :

$$\frac{dt}{2\nu} \|u^0\|_{L^\infty(\Omega)}^2 \leq 1. \quad (3.35)$$

And the L^2 CFL condition is :

$$\frac{dt(dx)^2}{2\nu} \|u^0\|_{L^2(\Omega)}^2 \leq 1. \quad (3.36)$$

List of Tables

1	Kolmogorov N-width of some partial differential equations.	6
2	Error metrics in the $\ \cdot\ _X$ norm.	23
3	Error metrics of the GPR output $h(\alpha(\mu))$ with respect to the target value $\beta(\mu)$ for different kernels and $\mu \in \mathbb{P}_{h,\text{test}}$	23
4	Computation time per sample of different methods for orthogonal modes.	24
5	Error metrics of the EIM for different test cases in the $\ \cdot\ _X$ norm.	26
6	Error metrics of the GPR output $h(\alpha(\mu))$ with respect to the target value $\beta(\mu)$ for different kernels and $\mu \in \mathbb{P}_{h,\text{test}}$	26
7	Comparison of different methods for advection.	31

List of Figures

1	Architecture of the different proposed methods of assimilation.	17
2	Modes of the numerical POD and their $\ \cdot\ _X^2$ -error with respect to the theoretical mode.	21
3	Energy histogram per modes (left). Cumulative energy $I = \sum_{i \leq n} \mathcal{E}_i / \sum_i \mathcal{E}_i$ (right).	22
4	Comparison of the numerical coefficients with the theory of the orthogonal test case.	22
5	Trajectory $u^\circ(x; \mu)$ for a given $\mu = 0.5$, the observed function y° and the sensors location.	22
6	Relative error $\ \mathbf{u}^\circ - \mathbf{u}_{\text{EIM}}\ _X / \ \mathbf{u}^\circ\ _X$ over $\mathbb{P}_{h,\text{test}}$ for different value of the noise scale σ when the number of sensors m varies and $n_{\text{EIM}} = 2$ if $m > 1$, $n_{\text{EIM}} = 1$ if $m = 1$	23
7	Error comparison for $\mu \in \mathbb{P}_{h,\text{test}}$, $m = N_{\text{modes}} = 2$, $\sigma = 10^{-1}$	24
8	Performance comparison of different methods when the number of sensors $m = 2^i$, $i = 0, \dots, 12$ and $n = \bar{n} = 1$	25
9	Relative error of reconstruction $\ \mathbf{u}_{n_d}(\mu) - \mathbf{u}^\star(\mu)\ _X / \ \mathbf{u}_{n_d}\ _X$ depending on the noise level σ when $m = 2^i$, $i = 1, \dots, 12$ and $n = \bar{n} = 1$	25
10	Energy histogram per modes (left). Cumulative energy $I = \sum_{i \leq n} \mathcal{E}_i / \sum_i \mathcal{E}_i$ (right).	26
11	Error on reconstructing $\mathbf{u}_{n_d}$ with a GPR trained on $n = 2, \bar{n} = 125$	27
12	Error on reconstructing $\mathbf{u}_{n_d}$ with a GPR trained on $n = 2, \bar{n} = 150$	27
13	Different learning error metrics for a GPR trained with $n = 2, \bar{n} = 125$	28
14	The first trajectory of the dataset, the first mode of the POD basis and the 100-th mode of the POD basis.	28
15	$\beta_i(\alpha_1)$ plotted as function of $\alpha_1(c)$ when the parameter $c \in \mathbb{P}$ varies and $i = 1, 50, 100$	29
16	Coefficient $\beta_{100}(\alpha_1, \alpha_2)$ plotted for the different values of $\alpha_1(c), \alpha_2(c)$ when the parameter $c \in \mathbb{P}$ varies.	29
17	Relative error of the GPR over the samples $\mu \in \mathbb{P}_{h,\text{test}}$. Mean relative $\mathbf{L}^2$ -error : 5.98×10^{-05}	30
18	Error comparison for $\mu \in \mathbb{P}_{h,\text{test}}$, $m = N_{\text{modes}} = 2$, $\sigma = 10^{-1}$. Advection case.	30
19	$\mathbf{L}^2$ -error for each samples of the GPR, EIM and Picard data assimilation method.	31
20	Graphic analysis of the GPR sensitivity to noise.	31

References

- [AG21] ERN A. and J-L. GUERMOND. *Finite Elements I: Approximation and Interpolation*. Springer International Publishing, 2021. DOI: [10.1007/978-3-030-56341-7](https://doi.org/10.1007/978-3-030-56341-7).
- [AGH+26] J. AGHILI et al. “Nonlinear compressive reduced basis approximation : when Taylor meets Kolmogorov”. In: (2026). DOI: [10.48550/ARXIV.2601.13712](https://doi.org/10.48550/ARXIV.2601.13712).
- [AGU23] F. ARBES, C. GREIF, and K. URBAN. *The Kolmogorov N-width for linear transport: Exact representation and the influence of the data*. 2023. DOI: [10.48550/ARXIV.2305.00066](https://doi.org/10.48550/ARXIV.2305.00066).
- [ARE+26] S. ARES DE PARGA et al. “Nonlinear projection-based model order reduction with machine learning regression for closure error modeling in the latent space”. In: *Computer Methods in Applied Mechanics and Engineering* 448 (2026). DOI: [10.1016/j.cma.2025.118443](https://doi.org/10.1016/j.cma.2025.118443).
- [BAR+04] M. BARRAULT et al. “An ‘empirical interpolation’ method: application to efficient reduced-basis discretization of partial differential equations”. In: *Comptes Rendus. Mathématique* 339 (2004). ISSN: 1778-3569. DOI: [10.1016/j.crma.2004.08.006](https://doi.org/10.1016/j.crma.2004.08.006).
- [BB88] JONATHAN BARZILAI and JONATHAN M. BORWEIN. “Two-Point Step Size Gradient Methods”. In: *IMA Journal of Numerical Analysis* (1988). ISSN: 1464-3642. DOI: [10.1093/imanum/8.1.141](https://doi.org/10.1093/imanum/8.1.141).
- [BAY+18] A. G. BAYDIN et al. “Automatic Differentiation in Machine Learning: a Survey”. In: *Journal of Machine Learning Research* 18.153 (2018), pp. 1–43. URL: <http://jmlr.org/papers/v18/17-468.html>.
- [BNS25] A. BENSALAH, A. NOUY, and J. SOFFO. *Nonlinear manifold approximation using compositional polynomial networks*. 2025. DOI: [10.48550/ARXIV.2502.05088](https://doi.org/10.48550/ARXIV.2502.05088).
- [BHK20] A. BLUM, J. HOPCROFT, and R. KANNAN. *Foundations of Data Science*. Cambridge University Press, 2020.
- [BRE11] H. BREZIS. *Functional Analysis, Sobolev Spaces and Partial Differential Equations*. Springer New York, 2011. DOI: [10.1007/978-0-387-70914-7](https://doi.org/10.1007/978-0-387-70914-7).
- [BUL12] V. BULJAK. *Inverse Analyses with Model Reduction*. Springer Berlin Heidelberg, 2012. DOI: [10.1007/978-3-642-22703-5](https://doi.org/10.1007/978-3-642-22703-5).
- [CEA64] J. CEA. “Approximation variationnelle des problèmes aux limites”. In: *Annales de l’Institut Fourier* (1964). DOI: [10.5802/aif.181](https://doi.org/10.5802/aif.181).
- [COL12] R. COLEMAN. *Calculus on Normed Vector Spaces*. Springer New York, 2012. DOI: [10.1007/978-1-4614-3894-6](https://doi.org/10.1007/978-1-4614-3894-6).
- [DE12] Daniele Antonio DI PIETRO and Alexandre ERN. *Mathematical Aspects of Discontinuous Galerkin Methods*. Springer, 2012. DOI: [10.1007/978-3-642-22980-0](https://doi.org/10.1007/978-3-642-22980-0).
- [EG04] A. ERN and J-L. GUERMOND. *Theory and Practice of Finite Elements*. 2004. DOI: [10.1007/978-1-4757-4355-5](https://doi.org/10.1007/978-1-4757-4355-5).
- [FIL24] F. FILBERT. *Analyse numérique : algorithme et étude mathématique*. 2nd edition. Dunod, 2024. ISBN: 9782100873791.
- [GON+19] H. GONG et al. “PBDW method for state estimation: error analysis for noisy data and nonlinear formulation”. In: 2019. DOI: <https://doi.org/10.48550/arXiv.1906.00810>.
- [GU19] C. GREIF and K. URBAN. *Decay of the Kolmogorov N-width for wave problems*. 2019. DOI: [10.48550/ARXIV.1903.08488](https://doi.org/10.48550/ARXIV.1903.08488).
- [HRS16] J. S. HESTHAVEN, G. ROZZA, and B. STAMM. *Certified Reduced Basis Methods for Parametrized Partial Differential Equations*. Springer International Publishing, 2016. DOI: [10.1007/978-3-319-22470-1](https://doi.org/10.1007/978-3-319-22470-1).

- [HJ12] R. A. HORN and C. R. JOHNSON. *Matrix Analysis*. Cambridge University Press, 2012. DOI: [10.1017/cbo9781139020411](https://doi.org/10.1017/cbo9781139020411).
- [JAI+17] S. JAIN et al. “A quadratic manifold for model order reduction of nonlinear structural dynamics”. In: *Computers & Structures* 188 (Aug. 2017), pp. 80–94. ISSN: 0045-7949. DOI: [10.1016/j.compstruc.2017.04.005](https://doi.org/10.1016/j.compstruc.2017.04.005).
- [KOL36] A. KOLMOGOROFF. “Über die beste Annäherung von Funktionen einer gegebenen Funktionenklasse”. In: *Annals of Mathematics* 37 (1936), pp. 107–110.
- [KYT96] F. K. KYTHE. *Fundamental Solutions for Differential Operators and Applications*. Birkhäuser Boston, 1996. DOI: [10.1007/978-1-4612-4106-5](https://doi.org/10.1007/978-1-4612-4106-5).
- [LAN96] C. LANCZOS. *Linear Differential Operators*. Society for Industrial and Applied Mathematics, 1996. DOI: [10.1137/1.9781611971187](https://doi.org/10.1137/1.9781611971187).
- [LE 13] H. LE DRET. *Équations aux dérivées partielles elliptiques non linéaires*. Springer Berlin Heidelberg, 2013. DOI: [10.1007/978-3-642-36175-3](https://doi.org/10.1007/978-3-642-36175-3).
- [LIA+02] Y.C. LIANG et al. “Proper Orthogonal Decomposition and Its Applications—Part I: Theory”. In: *Journal of Sound and Vibration* (2002). DOI: [10.1006/jsvi.2001.4041](https://doi.org/10.1006/jsvi.2001.4041).
- [MPT02] Y. MADAY, A. T. PATERA, and G. TURINICI. “Global a priori convergence theory for reduced-basis approximations of single-parameter symmetric coercive elliptic partial differential equations”. In: *Comptes Rendus. Mathématique* (2002). DOI: [10.1016/s1631-073x\(02\)02466-4](https://doi.org/10.1016/s1631-073x(02)02466-4).
- [MAD+15] Y. MADAY et al. “PBDW State Estimation: Noisy Observations; Configuration-Adaptive Background Spaces; Physical Interpretations”. In: *ESAIM: Proceedings and Surveys* (2015). DOI: [10.1051/proc/201550008](https://doi.org/10.1051/proc/201550008).
- [MM13] Yvon Maday and Olga Mula. “A Generalized Empirical Interpolation Method: Application of Reduced Basis Techniques to Data Assimilation”. In: *Analysis and Numerics of Partial Differential Equations*. Springer Milan, 2013, pp. 221–235. ISBN: 9788847025929. DOI: [10.1007/978-88-470-2592-9_13](https://doi.org/10.1007/978-88-470-2592-9_13).
- [MS20] C. MARCATI and C. SCHWAB. “Analytic regularity for the incompressible Navier-Stokes equations in polygons”. In: (2020). DOI: [10.48550/ARXIV.2004.11264](https://doi.org/10.48550/ARXIV.2004.11264).
- [MEN90] B MENDELSON. *Introduction to Topology*. 3rd ed. Dover Publications, 1990. ISBN: 978-0486663524.
- [NEI10] R. D. NEIDINGER. “Introduction to Automatic Differentiation and MATLAB Object-Oriented Programming”. In: *SIAM Review* 52 (2010). DOI: [10.1137/080743627](https://doi.org/10.1137/080743627).
- [NIU+25] T. NIU et al. “Machine Learning-Based Bias Correction for the GEIM Model”. In: *Numerical Mathematics: Theory, Methods and Applications* (Oct. 2025). ISSN: 1004-8979. DOI: [10.4208/nmtma.oa-2025-0025](https://doi.org/10.4208/nmtma.oa-2025-0025).
- [OS16] M. OHLBERGER and RAVE S. “Reduced Basis Methods: Success, Limitations and Future Challenges”. In: (2016). v2. URL: <https://arxiv.org/abs/1511.02021v2>.
- [PD20] B. PEHERSTORFER and S. DRMAC Z. and GUGERCIN. “Stability of Discrete Empirical Interpolation and Gappy Proper Orthogonal Decomposition with Randomized and Deterministic Sampling Points”. In: *SIAM Journal on Scientific Computing* (2020). DOI: [10.1137/19m1307391](https://doi.org/10.1137/19m1307391).
- [QMF16] A. QUARTERONI, A. MANZONI, and N. FEDERICO. *Reduced Basis Methods for Partial Differential Equations: An Introduction*. MS&A – Modeling, Simulation and Applications. Springer, 2016. DOI: [10.1007/978-3-319-15431-2](https://doi.org/10.1007/978-3-319-15431-2).
- [QSS06] A. QUARTERONI, R. SACCO, and F. SALERI. *Numerical Mathematics*. 2nd ed. 2006. DOI: [10.1007/b98885](https://doi.org/10.1007/b98885).

- [RW05] C. E. RASMUSSEN and C. K. I. WILLIAMS. *Gaussian Processes for Machine Learning*. The MIT Press, 2005. DOI: [10.7551/mitpress/3206.001.0001](https://doi.org/10.7551/mitpress/3206.001.0001).
- [RAU91] J. RAUCH. *Partial Differential Equations*. Springer New York, 1991. DOI: [10.1007/978-1-4612-0953-9](https://doi.org/10.1007/978-1-4612-0953-9).
- [RHE88] W. C. RHEINBOLDT. “On the computation of multi-dimensional solution manifolds of parametrized equations”. In: *Numerische Mathematik* (1988). DOI: [10.1007/bf01395883](https://doi.org/10.1007/bf01395883).
- [SP24] P. SCHWERDTNER and B. PEHERSTORFER. *Greedy construction of quadratic manifolds for nonlinear dimensionality reduction and nonlinear model reduction*. 2024. DOI: [10.48550/ARXIV.2403.06732](https://doi.org/10.48550/ARXIV.2403.06732).
- [SEN90] B. SENDOV. *Hausdorff Approximations*. Ed. by Gerald Beer. Springer Netherlands, 1990. DOI: [10.1007/978-94-009-0673-0](https://doi.org/10.1007/978-94-009-0673-0).
- [TAY91] M. E. TAYLOR. “Pseudodifferential Operators and Nonlinear PDE”. In: (1991). DOI: [10.1007/978-1-4612-0431-2](https://doi.org/10.1007/978-1-4612-0431-2).