



Covariance-Informed Subspace: a Gradient-free Dimension Reduction for Adaptive Bayesian Inference

MASCOTNUM2025 - PhD Day

PhD student: Nadège Polette^{1,2},

Supervisors:

Alexandrine Gesret¹, Pierre Sochala², Olivier Le Maître³



Table of contents

Context

Field parametrization

Dimension reduction

Applications



1. Context: Detection and analysis of seismic events

Global scale

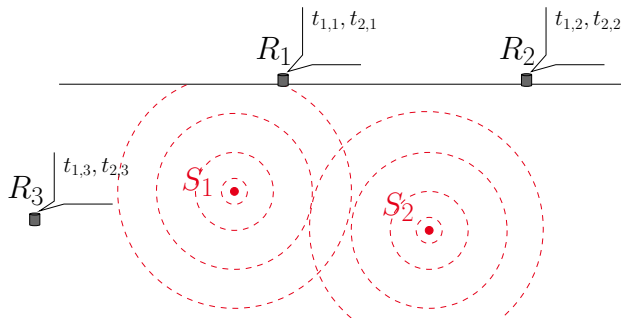
- International treaties (CTBT, NTP)
- Environment monitoring (IMS)

Regional scale

- Tsunami and earthquake alerts
- Risk prevention

Local scale

- Subsurface knowledge
- Exploitation



1. Context: Detection and analysis of seismic events

Global scale

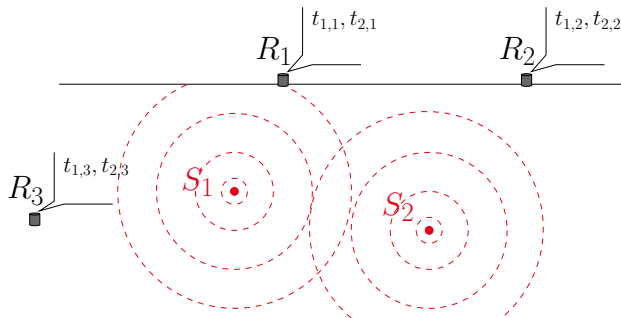
- International treaties (CTBT, NTP)
- Environment monitoring (IMS)

Regional scale

- Tsunami and earthquake alerts
- Risk prevention

Local scale

- Subsurface knowledge
- Exploitation



$$F(S) = d$$

F : forward model

S : source parameters

d : data

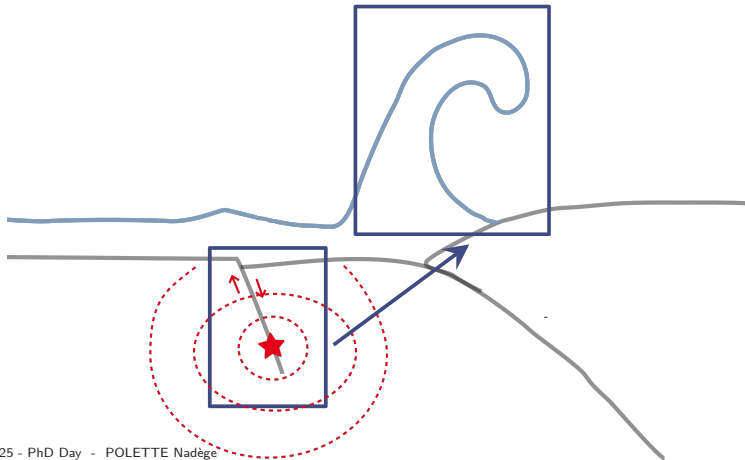
Objective - retrieve S from d

- fast
- accurately
- with uncertainties



1. Context: Uncertainty on physical parameters

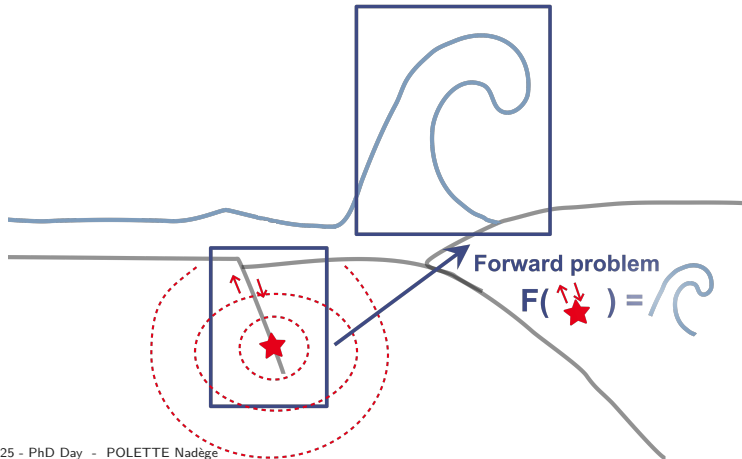
- Earthquakes could cause tsunamis within minutes/hours





1. Context: Uncertainty on physical parameters

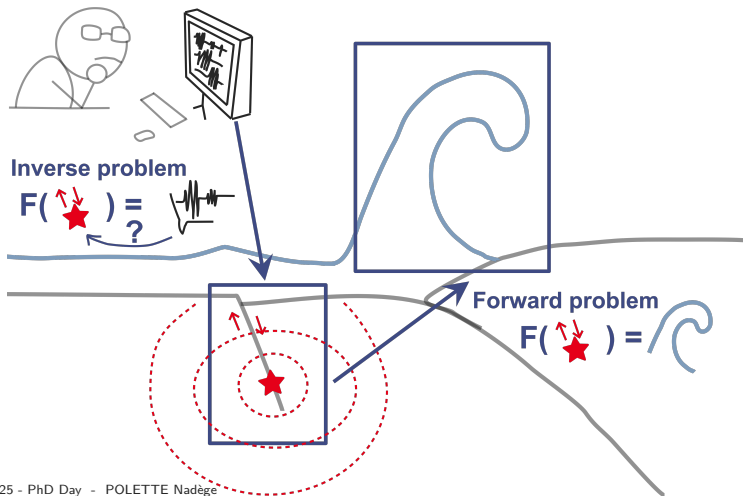
- Earthquakes could cause tsunamis within minutes/hours
- Source parameters → Simulation





1. Context: Uncertainty on physical parameters

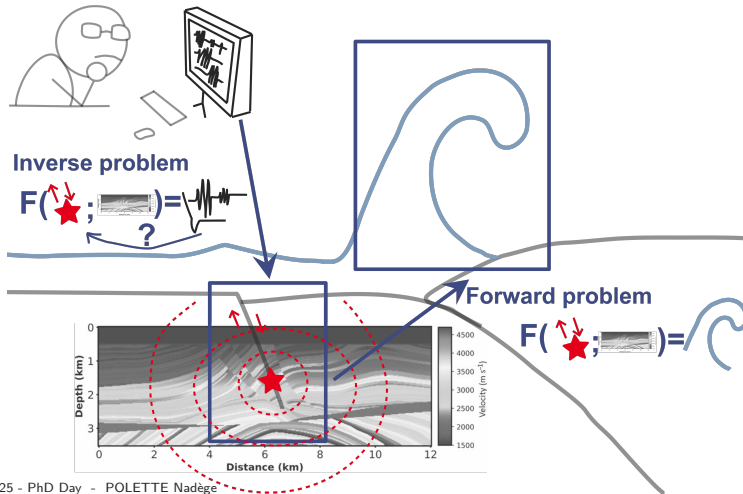
- Earthquakes could cause tsunamis within minutes/hours
- Data → Source parameters → Simulation





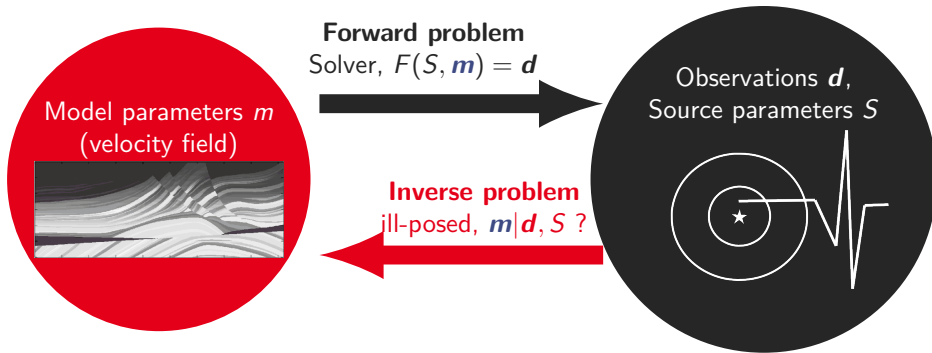
1. Context: Uncertainty on physical parameters

- Earthquakes could cause tsunamis within minutes/hours
- Data $\rightarrow$ Source parameters $\rightarrow$ Simulation
- Poorly known parameters, e.g. the velocity field





1. Context: Inverse problem

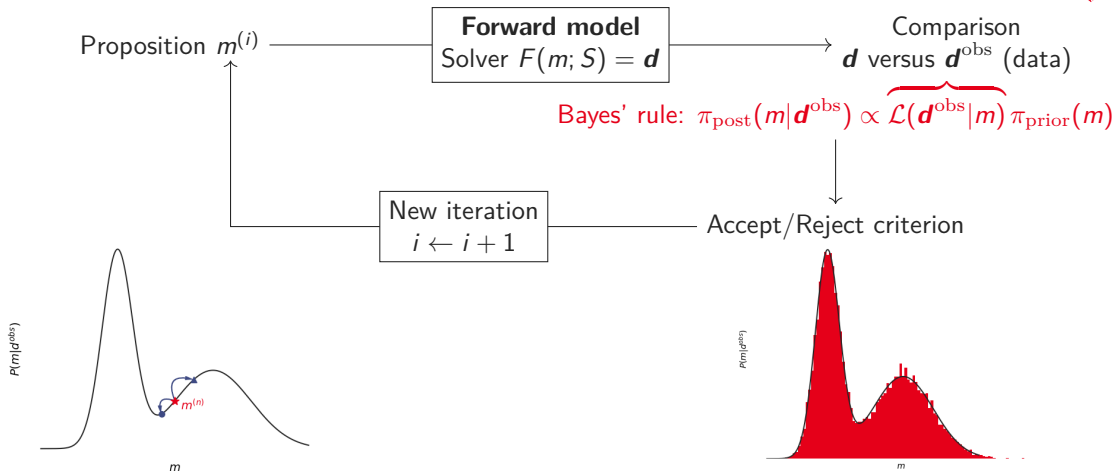


Objective: to characterize the velocity field m and its uncertainty from indirect observations d
⇒ to find the probability distribution of the field knowing the observations $\pi_{\text{post}}(m|d^{\text{obs}})$

Tarantola, *SIAM*, 2005; Noble et al., *GJI*, 2014



1. Context: Bayesian inference and Markov chain Monte Carlo



Sivia and Skilling, *Oxford*, 2006;

Doucet et al., *Springer NY*, 2013



1. Context: Bayesian inference and Markov chain Monte Carlo

Proposition

m infinite dimensional
→ dimension reduction

$m^{(i)}$

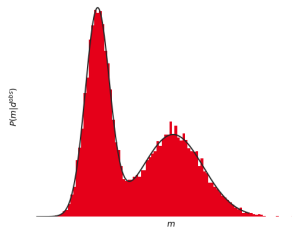
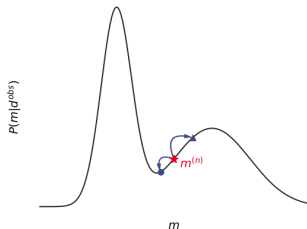
Forward model
Solver $F(m; S) = d$

Comparison
 d versus d^{obs} (data)

Bayes' rule: $\pi_{\text{post}}(m|d^{\text{obs}}) \propto \mathcal{L}(d^{\text{obs}}|m) \pi_{\text{prior}}(m)$

New iteration
 $i \leftarrow i + 1$

Accept/Reject criterion



Sivia and Skilling, Oxford, 2006;

Doucet et al., Springer NY, 2013

Table of contents

Context

Field parametrization

Dimension reduction

Applications

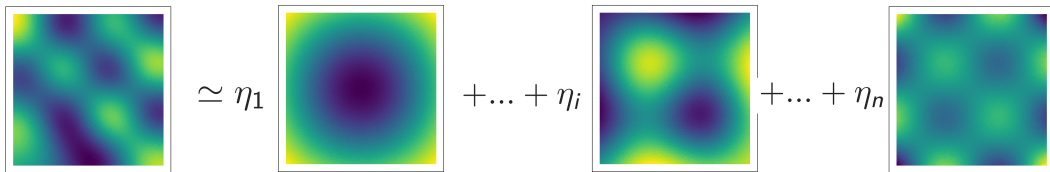


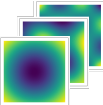


2. Field parametrization: Modal representation

Assumption - $m \sim \mathcal{N}(0, K)$

Karhunen–Loève decomposition - $m(\mathbf{x}) = \sum_{i=1}^n \sqrt{\lambda_i} u_i(\mathbf{x}) \eta_i, \mathbf{x} \in \Omega$



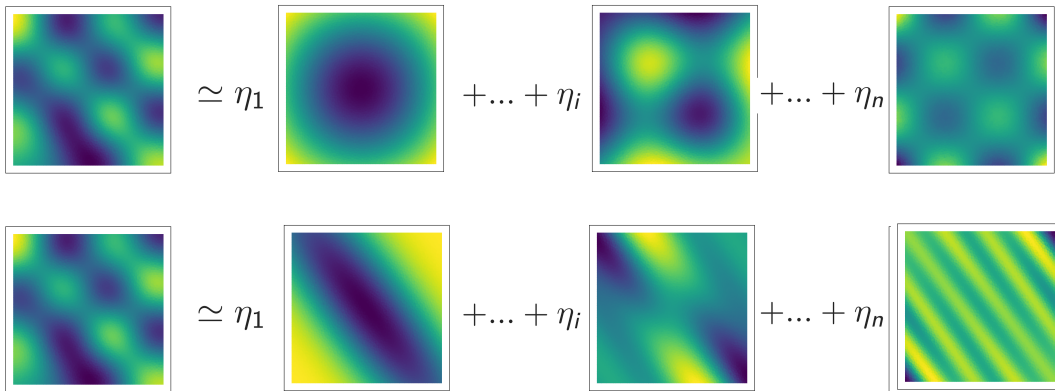
■  obtained using eigendecomposition of K : $\int_{\Omega} K(\mathbf{x}, \mathbf{y}) u_i(\mathbf{x}) d\mathbf{x} = \lambda_i u_i(\mathbf{y})$

■ $\pi_{\text{post}}(m | \mathbf{d}^{\text{obs}}) \Rightarrow \pi_{\text{post}}(\boldsymbol{\eta} | \mathbf{d}^{\text{obs}}) \propto \mathcal{L}(\mathbf{d}^{\text{obs}} | \boldsymbol{\eta}) \pi(\boldsymbol{\eta}), \quad \boldsymbol{\eta} \sim \mathcal{N}(0, \mathbf{I}_n)$



2. Field parametrization: Kernel choice

Problem - the choice of K is not always trivial



$$\blacksquare \quad \pi_{\text{post}}(m|\mathbf{d}^{\text{obs}}) \Rightarrow \pi_{\text{post}}(\boldsymbol{\eta}, \mathbf{q}|\mathbf{d}^{\text{obs}}) \propto \mathcal{L}(\mathbf{d}^{\text{obs}}|\boldsymbol{\eta}, \mathbf{q})\pi(\boldsymbol{\eta})\pi(\mathbf{q}), \quad \boldsymbol{\eta} \sim \mathcal{N}(0, \mathbf{I}_n), \quad \mathbf{q} \sim \pi_{\mathbf{q}}$$

Table of contents

Context

Field parametrization

Dimension reduction

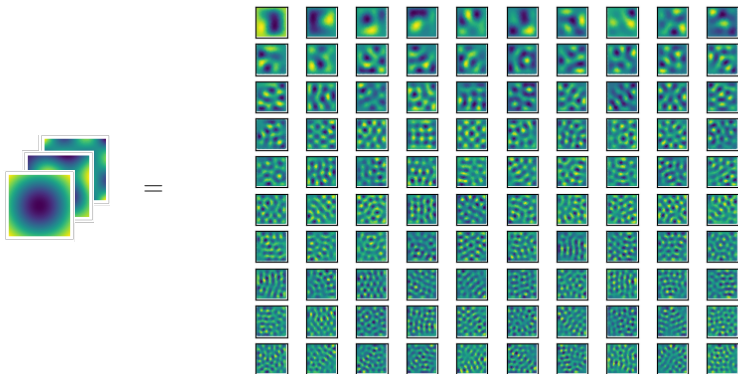
Applications





3. Dimension reduction: Context

Problem - Parametrisation may require a large number of modes

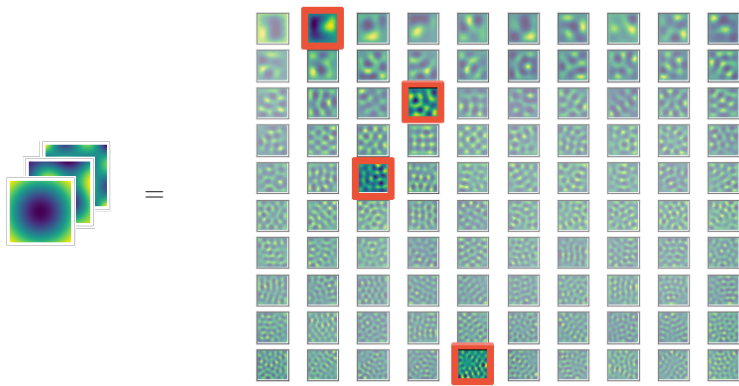


⇒ Computationally expensive



3. Dimension reduction: Context

Idea - Select only a sub-collection



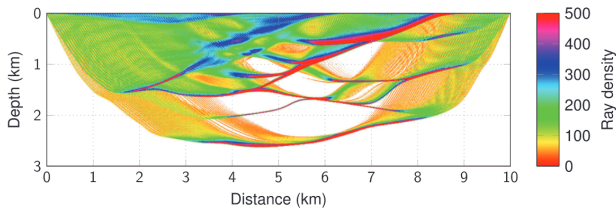
⇒ How to choose the selected modes ?



3. Dimension reduction: Informed subspace

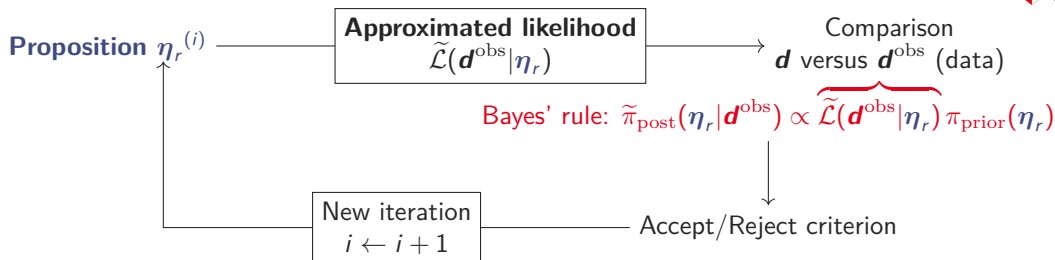
Assumption - the data are informative only on a low-dimensional subspace

$$\pi_{\text{post}}(\boldsymbol{\eta}|\mathbf{d}^{\text{obs}}) \simeq \tilde{\pi}_{\text{post}}(\boldsymbol{\eta}|\mathbf{d}^{\text{obs}}) \propto \tilde{\mathcal{L}}(\mathbf{d}^{\text{obs}}|\boldsymbol{\eta}_r)\pi(\boldsymbol{\eta}_r)\pi(\boldsymbol{\eta}_{\perp}|\boldsymbol{\eta}_r), \quad \boldsymbol{\eta} = A_r\boldsymbol{\eta}_r + A_{\perp}\boldsymbol{\eta}_{\perp}$$



Data limitations: ray density map (Luu et al., *Geoph. Prosp.*, 2020)

3. Dimension reduction: Informed subspace



- $\eta_{\perp}$ is sampled according to $\pi_{\text{prior}}(\eta_{\perp}|\eta_r)$
- **Approximated likelihood**

$$\tilde{\mathcal{L}}(\mathbf{d}^{\text{obs}}|\eta_r) = \mathbb{E}_{\pi_{\text{prior}}}(\mathcal{L}|\eta_r) = \int_{\Xi_{\perp}} \mathcal{L}(\mathbf{d}^{\text{obs}}|\eta) \pi_{\text{prior}}(\eta_{\perp}|\eta_r) d\eta_{\perp}$$

- optimal with respect to (i) the $L^2\pi$ -norm, (ii) the Kullback–Leibler divergence [Zahm2022, Section 2], and more generally (iii) all expected Bregman divergences [Banerjee2005]



3. Dimension reduction: Informed subspace

Assumption - the data are informative only on a low-dimensional subspace

$$\pi_{\text{post}}(\boldsymbol{\eta}|\mathbf{d}^{\text{obs}}) \simeq \tilde{\pi}_{\text{post}}(\boldsymbol{\eta}|\mathbf{d}^{\text{obs}}) \propto \tilde{\mathcal{L}}(\mathbf{d}^{\text{obs}}|\boldsymbol{\eta}_r)\pi(\boldsymbol{\eta}_r)\pi(\boldsymbol{\eta}_{\perp}|\boldsymbol{\eta}_r), \quad \boldsymbol{\eta} = A_r\boldsymbol{\eta}_r + A_{\perp}\boldsymbol{\eta}_{\perp}$$

⇒ State-of-the-art methods are *gradient-based*.

- $H = \int_{\Xi} \nabla \log \mathcal{L}(\mathbf{d}^{\text{obs}}|\boldsymbol{\eta}) \left(\nabla \log \mathcal{L}(\mathbf{d}^{\text{obs}}|\boldsymbol{\eta}) \right)^{\top} \nu(\boldsymbol{\eta}) d\boldsymbol{\eta}$
- The reduction relies on the dominant eigenspace of the pencil $(H, C_{\text{prior}}^{-1} = I_n)$:

$$[u_1 \dots u_r], \quad \lambda_1 \geq \dots \geq \lambda_n, \quad Hu_i = \lambda_i C_{\text{prior}}^{-1} u_i$$



3. Dimension reduction: Informed subspace

Assumption - the data are informative only on a low-dimensional subspace

$$\pi_{\text{post}}(\boldsymbol{\eta}|\mathbf{d}^{\text{obs}}) \simeq \tilde{\pi}_{\text{post}}(\boldsymbol{\eta}|\mathbf{d}^{\text{obs}}) \propto \tilde{\mathcal{L}}(\mathbf{d}^{\text{obs}}|\boldsymbol{\eta}_r)\pi(\boldsymbol{\eta}_r)\pi(\boldsymbol{\eta}_{\perp}|\boldsymbol{\eta}_r), \quad \boldsymbol{\eta} = A_r\boldsymbol{\eta}_r + A_{\perp}\boldsymbol{\eta}_{\perp}$$

⇒ State-of-the-art methods are *gradient-based*.

- $H = \int_{\Xi} \nabla \log \mathcal{L}(\mathbf{d}^{\text{obs}}|\boldsymbol{\eta}) \left(\nabla \log \mathcal{L}(\mathbf{d}^{\text{obs}}|\boldsymbol{\eta}) \right)^{\top} \nu(\boldsymbol{\eta}) d\boldsymbol{\eta}$
- The reduction relies on the dominant eigenspace of the pencil $(H, C_{\text{prior}}^{-1} = I_n)$:

$$[u_1 \dots u_r], \quad \lambda_1 \geq \dots \geq \lambda_n, \quad Hu_i = \lambda_i C_{\text{prior}}^{-1} u_i$$

⇒ Gradients can be expensive/unavailable



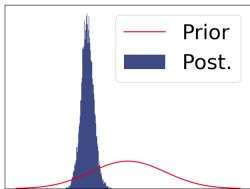
3. Dimension reduction: Gradient-free formulation

Bayesian linear Gaussian case - Linear forward model, Gaussian posterior

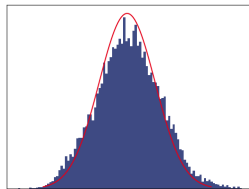
Spantini's corollary - The reduction equivalently relies on the minor eigenspace of $(C_{\text{post}}, C_{\text{prior}} = I_n)$

$$[v_1 \dots v_r], \quad \nu_1 \leq \dots \leq \nu_n, \quad C_{\text{post}} v_i = \nu_i C_{\text{prior}} v_i$$

New corollary - In the (BLG) case, considering an approximation of the form $\tilde{\pi}_{\text{post}}$, this projector is also optimal to approximate C_{post} .



(a) Informed direction



(b) Non-informed direction



3. Dimension reduction: Gradient-free formulation

Posterior covariance approximation

- C_{post} is unknown
- Inputs: m samples $\boldsymbol{\eta}^{(i)} \sim \tilde{\pi}_{\text{post}}, \mathcal{L}(\boldsymbol{\eta}^{(i)}), \tilde{\mathcal{L}}(\boldsymbol{\eta}_r^{(i)})$
- Weights: $\omega^{(i)} = \mathcal{L}(\boldsymbol{\eta}^{(i)}) / \tilde{\mathcal{L}}(\boldsymbol{\eta}_r^{(i)})$
- Weighted Monte Carlo estimator

$$\hat{C}_{\text{post}} = \frac{\sum_{i=1}^m \omega^{(i)}}{\left(\sum_{i=1}^m \omega^{(i)}\right)^2 - \sum_{i=1}^m (\omega^{(i)})^2} \sum_{i=1}^m \omega^{(i)} (\boldsymbol{\eta}^{(i)} - \bar{\boldsymbol{\eta}})(\boldsymbol{\eta}^{(i)} - \bar{\boldsymbol{\eta}})^\top, \text{ with } \bar{\boldsymbol{\eta}} = \frac{\sum_{i=1}^m \omega^{(i)} \boldsymbol{\eta}^{(i)}}{\sum_{i=1}^m \omega^{(i)}}$$

- Tempering/shrinkage etc. available



3. Dimension reduction: Gradient-free formulation

Extension to non-linear cases - A bound can be derived from (Cui and Tong, *Bern.*, 2022)

Bound of the final approximation quality

The Hellinger distance between the true posterior P and the numerical approximated posterior $P_{\Pi_r}^{(N)}$ is bounded with

$$D_H(P, P_{\Pi_r}^{(N)}) \leq \sqrt{\mathbb{E}_{P_{\Pi_r}} \left(\text{Var}_{\pi_{\perp|r}}(X_{\perp}|X_r)(\sqrt{\omega(X)}) \right)} + \sqrt{\frac{2}{N}} \sqrt{\mathbb{E}_{P_{\Pi_r}} \left(\text{Var}_{\pi_{\perp|r}}(X_{\perp}|X_r)(\omega(X)) \right)}.$$

Table of contents

Context

Field parametrization

Dimension reduction

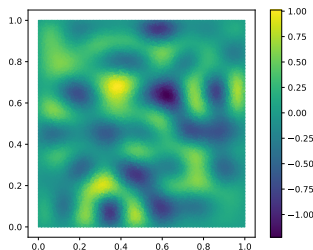
Applications



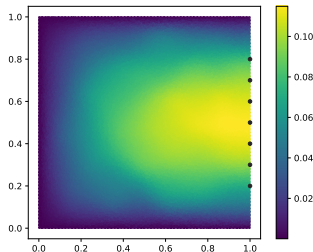
4. Application: Steady-state diffusion equation

$$\begin{cases} -\nabla (\kappa(x) \cdot \nabla u(x)) = 1, & x \in \Omega = [0, 1]^2, \\ u(x) = 0, & x \in \Gamma_D = \{x, \quad x_1 = 0 \cup x_2 = 0 \cup x_2 = 1\}, \\ \nabla u(x) \cdot \mathbf{n} = 0, & x \in \Gamma_N = \{x, \quad x_1 = 1\}. \end{cases}$$

$\log \kappa \sim \mathcal{N}(0, k)$, k an exponential kernel with correlation length 0.02



(a) True log-field

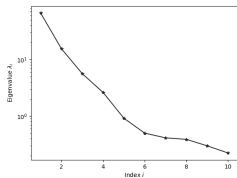


(b) True solution

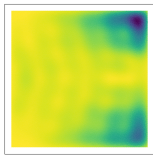
Constantine et al., *SIAM JSC*, 2016



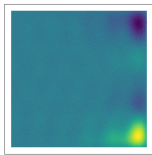
4. Application: Steady-state diffusion equation



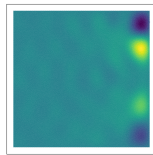
(a) Eigenvalues



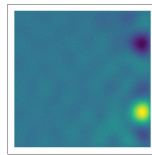
(b) First direction



(c) Second direction



(d) Third direction

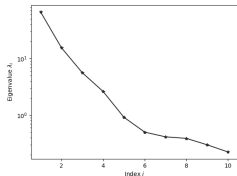


(e) Fourth direction

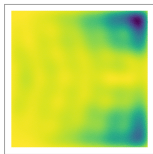
$(H, C_{\text{prior}}^{-1}), n = 1,000$, prior weighted averaged



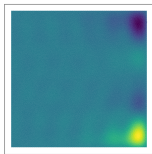
4. Application: Steady-state diffusion equation



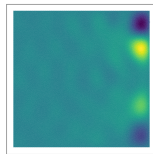
(a) Eigenvalues



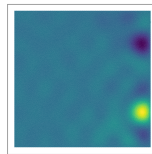
(b) First direction



(c) Second direction

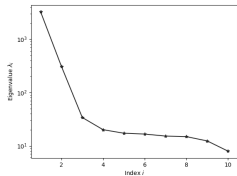


(d) Third direction

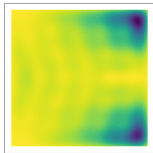


(e) Fourth direction

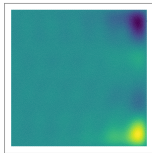
$(H, C_{\text{prior}}^{-1}), n = 1,000$, prior weighted averaged



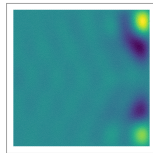
(a) Eigenvalues



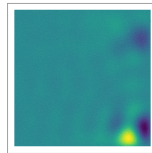
(b) First direction



(c) Second direction



(d) Third direction

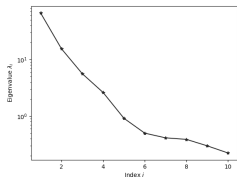


(e) Fourth direction

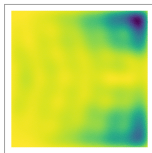
$(H, C_{\text{prior}}^{-1}), n = 1,000$, prior averaged



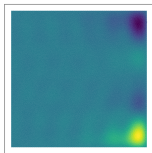
4. Application: Steady-state diffusion equation



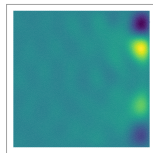
(a) Eigenvalues



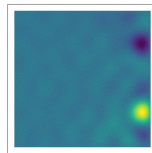
(b) First direction



(c) Second direction

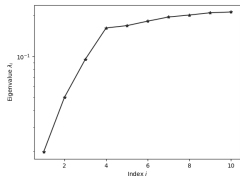


(d) Third direction

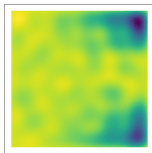


(e) Fourth direction

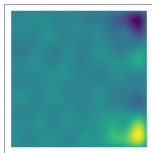
$(H, C_{\text{prior}}^{-1}), n = 1,000$, prior weighted averaged



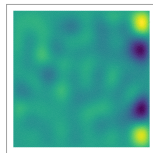
(a) Eigenvalues



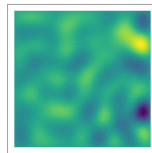
(b) First direction



(c) Second direction



(d) Third direction

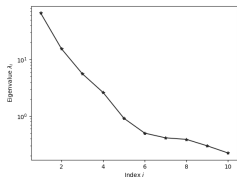


(e) Fourth direction

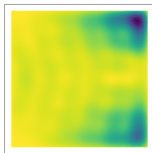
$(\hat{C}_{\text{post}}, C_{\text{prior}}), n = 10,000$



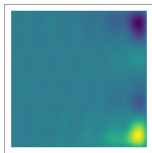
4. Application: Steady-state diffusion equation



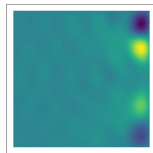
(a) Eigenvalues



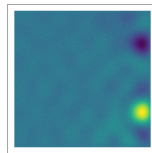
(b) First direction



(c) Second direction

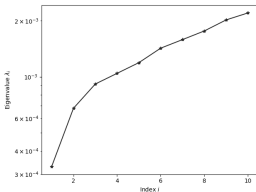


(d) Third direction

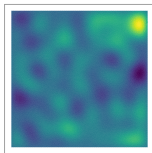


(e) Fourth direction

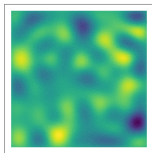
$(H, C_{\text{prior}}^{-1}), n = 1,000$, prior weighted averaged



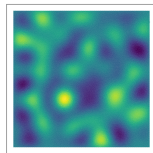
(a) Eigenvalues



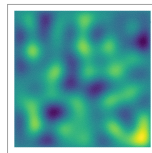
(b) First direction



(c) Second direction



(d) Third direction

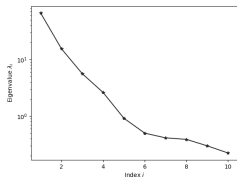


(e) Fourth direction

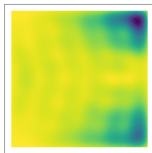
$(\hat{C}_{\text{post}}, C_{\text{prior}}), n = 1,000$



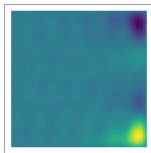
4. Application: Steady-state diffusion equation



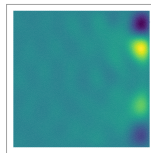
(a) Eigenvalues



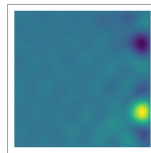
(b) First direction



(c) Second direction

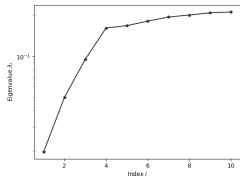


(d) Third direction

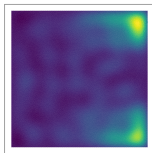


(e) Fourth direction

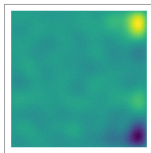
$(H, C_{\text{prior}}^{-1}), n = 1,000$, prior weighted averaged



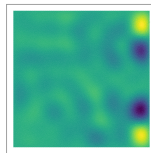
(a) Eigenvalues



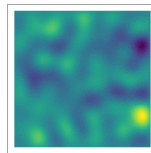
(b) First direction



(c) Second direction



(d) Third direction



(e) Fourth direction

$(\hat{C}_{\text{post}}, C_{\text{prior}}), n = 1,000$, 8 iterations.



4. Application: Bivariate cases

- Results for bivariate cases: $X = (x_1, x_2)$, $x_1 \sim \pi_{\text{post}}$, $x_2 \sim \mathcal{N}(0, 1)$

π_{post}	Illus	λ Grad.-based	Cov.-based
$\mathcal{N}(0, 1)$		✓	✓
$\mathcal{N}(0.5, 0.1)$		✓	✓
$\mathcal{N}(0.5, 1)$		✓	✗
Bi-modal, std $\simeq 0.5$		✓	✓
Bi-modal, std $\simeq 1$		✓	✗
Bi-modal, std $\simeq 2$		✓	✗



4. Application: Bivariate cases

- Results for bivariate cases: $X = (x_1, x_2)$, $x_1 \sim \pi_{\text{post}}$, $x_2 \sim \mathcal{N}(0, 1)$
- Correlation coefficient (c.c.) values: $\mathbb{C}\text{or}(X, \log \mathcal{L})$, $\mathbb{C}\text{or}(X^2, \log \mathcal{L})$, Chatterjee coef.

π_{post}	Illus	λ Grad.-based	Cov.-based	lin. c.c.	quad. c.c.	Chatterjee
$\mathcal{N}(0, 1)$		✓	✓	N.A.	N.A.	N.A.
$\mathcal{N}(0.5, 0.1)$		✓	✓	✓	✓	✓
$\mathcal{N}(0.5, 1)$		✓	✗	✓	✗	✓
Bi-modal, std $\simeq 0.5$		✓	✓	✗	✓	✓
Bi-modal, std $\simeq 1$		✓	✗	✗	✗	✓
Bi-modal, std $\simeq 2$		✓	✗	✗	✓	✓



Conclusion

- *Objective*: develop numerical methods to help account for field uncertainties in inverse problems
- *Problem*: the parametrization could require a high number of parameters
- *Covariance-Informed Subspace*: adaptive gradient-free dimension reduction
- *Application*: Steady-state diffusion problem $\rightarrow$ from dimension 100 to ~ 4
- Work in progress: using correlation coefficients and partial least squares; application to Global Ozone Monitoring System (GOMOS)

Thank you !

nadege.polette@minesparis.psl.eu

Keywords: inverse problem, Bayesian inference, MCMC, dimension reduction, gradient free