

Introduction aux méthodes de réduction de modèle pour les EDP

Marie Billaud-Friess

Centrale Méditerranée, Aix-Marseille Université
Institut de Mathématiques de Marseille

Master Class - M2 ANADEAL

20 Juin 2025

1. Introduction

2. Equation d'advection-diffusion dépendant de paramètres

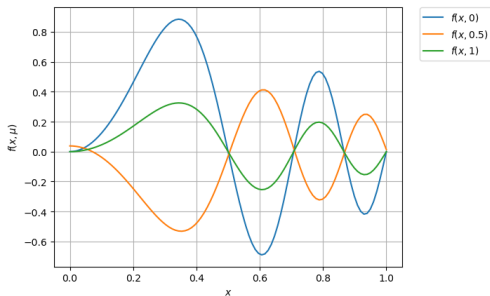
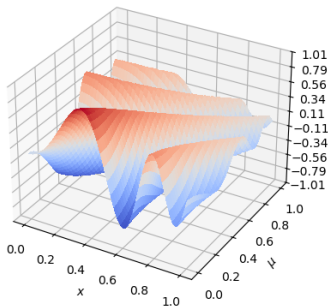
3. Equation de la chaleur

4. Un peu de théorie

5. Conclusion

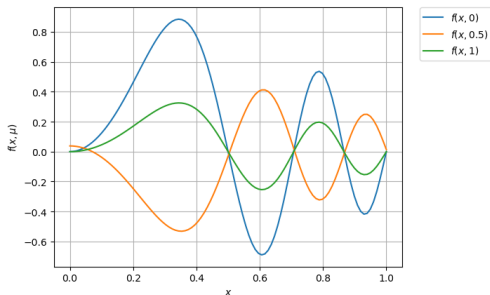
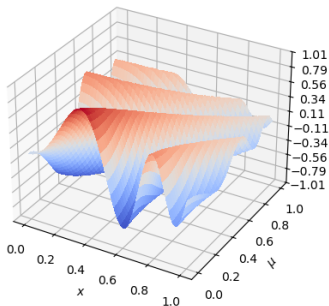
Considérons

$$f(x, \mu) = \exp(-x^2 - \mu) \sin(4\pi(x^2 - \mu^2)), (x, \mu) \in [0, 1]^2$$



Considérons

$$f(x, \mu) = \exp(-x^2 - \mu) \sin(4\pi(x^2 - \mu^2)), (x, \mu) \in [0, 1]^2$$



Question. Peut-on écrire la fonction "compliquée" f comme une somme de produits de fonctions "simples" dépendant de x et μ séparément ?

Réponse. Oui, après développement :

$$f(x, \mu) = \underbrace{\exp(-x^2) \sin(4\pi x^2)}_{\varphi_1(x)} \underbrace{\exp(-\mu) \cos(\mu^2)}_{\alpha_1(\mu)} + \underbrace{(-\exp(-x^2) \cos(4\pi x^2))}_{\varphi_2(x)} \underbrace{\exp(-\mu) \sin(\mu^2)}_{\alpha_2(\mu)},$$

c'est à dire

$$f(x, \mu) = \sum_{i=1}^2 \alpha_i(\mu) \varphi_i(x).$$

La fonction f est dite de rang 2.

Réponse. Oui, après développement :

$$f(x, \mu) = \underbrace{\exp(-x^2) \sin(4\pi x^2)}_{\varphi_1(x)} \underbrace{\exp(-\mu) \cos(\mu^2)}_{\alpha_1(\mu)} + \underbrace{(-\exp(-x^2) \cos(4\pi x^2))}_{\varphi_2(x)} \underbrace{\exp(-\mu) \sin(\mu^2)}_{\alpha_2(\mu)},$$

c'est à dire

$$f(x, \mu) = \sum_{i=1}^2 \alpha_i(\mu) \varphi_i(x).$$

La fonction f est dite de **rang 2**.

En pratique. Soient $\{x_i\}_{i=1}^n, \{\mu_i\}_{i=1}^m \subset [0, 1]$. Il s'agit de calculer

✓ seulement $2(m+n)$ évaluations de α_1, α_2 en $\{\mu_i\}_{i=1}^m$ et pour φ_1, φ_2 en $\{x_i\}_{i=1}^n$,

✗ contre nm évaluations de f en $\{(x_i, \mu_j)\}_{1 \leq i \leq n, 1 \leq j \leq m}$.

Approximation. Considérons

$$g(x, \mu) = \sum_{i \geq 1} 10^{-i} \exp(-x^2 - \mu) \sin(2^i \pi(x^2 - \mu^2)), (x, \mu) \in [0, 1]^2$$

Approximation. Considérons

$$g(x, \mu) = \sum_{i \geq 1} 10^{-i} \exp(-x^2 - \mu) \sin(2^i \pi(x^2 - \mu^2)), (x, \mu) \in [0, 1]^2$$

approchée par une fonction de rang $2k$

$$g_{2k}(x, \mu) = \sum_{i=1}^k 10^{-i} \exp(-x^2 - \mu) \sin(2^i \pi(x^2 - \mu^2)), (x, \mu) \in [0, 1]^2$$

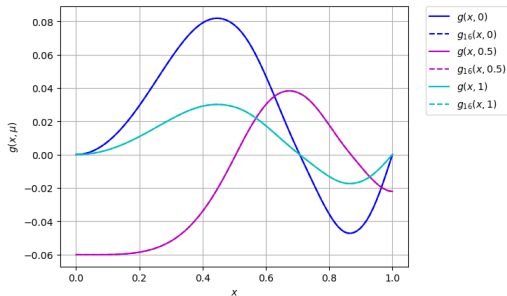
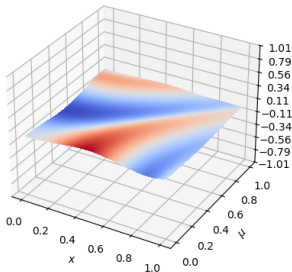
Exemple jouet ② : fonction approchable par une fonction à variables séparées

Approximation. Considérons

$$g(x, \mu) = \sum_{i \geq 1} 10^{-i} \exp(-x^2 - \mu) \sin(2^i \pi(x^2 - \mu^2)), (x, \mu) \in [0, 1]^2$$

approchée par une fonction de rang $2k$

$$g_{2k}(x, \mu) = \sum_{i=1}^k 10^{-i} \exp(-x^2 - \mu) \sin(2^i \pi(x^2 - \mu^2)), (x, \mu) \in [0, 1]^2$$



$k = 1$ rang = 2
 $k = 8$ rang = 16
 $k = 15$ rang = 30

$err_2 = 0.347220406696138$
 $err_2 = 3.616671153088143e - 08$
 $err_2 = 3.619439553869622e - 15$

① Certaines fonctions **admettent** une représentation **séparée de faible rang**.

ex. : La fonction f .

② D'autres peuvent **être approchées** par des fonctions qui admettent une représentation **séparée de faible rang**.

ex. : La fonction g .

↪ Une fonction sous format séparé de faible rang est **"moins" coûteuse à évaluer**.

Problèmes complexes issus de la physique, biologie, finance, chimie, ...

Problèmes complexes issus de la physique, biologie, finance, chimie, ...

1. formulation mathématique : équation aux dérivées partielles $\mathcal{F}(u) = 0$

Equation de la chaleur

$$\partial_t u - \Delta u = 0$$

Equation d'advection diffusion

$$-\kappa \Delta u + a \cdot \nabla u = f$$

Problèmes complexes issus de la physique, biologie, finance, chimie, ...

1. formulation mathématique : équation aux dérivées partielles $\mathcal{F}(u) = 0$

Equation de la chaleur

$$\partial_t u - \Delta u = 0$$

Equation d'advection diffusion

$$-\kappa \Delta u + a \cdot \nabla u = f$$

2. non résolvable analytiquement : résolution numérique $F_h(u_h) = 0$

Problèmes complexes issus de la physique, biologie, finance, chimie, ...

1. formulation mathématique : équation aux dérivées partielles $\mathcal{F}(u) = 0$

Equation de la chaleur

$$\partial_t u - \Delta u = 0$$

Equation d'advection diffusion

$$-\kappa \Delta u + a \cdot \nabla u = f$$

2. non résolvable analytiquement : résolution numérique $F_h(\mathbf{u}_h) = 0$
3. résolution numérique trop coûteuse : réduction de modèle $F_n(\mathbf{u}_n) = 0$

Problèmes complexes issus de la physique, biologie, finance, chimie, ...

1. formulation mathématique : **équation aux dérivées partielles** $\mathcal{F}(u) = 0$

Equation de la chaleur

$$\partial_t u - \Delta u = 0$$

Equation d'advection diffusion

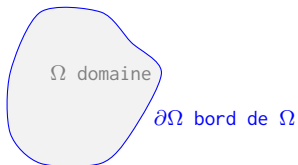
$$-\kappa \Delta u + a \cdot \nabla u = f$$

2. non résolvable analytiquement : **résolution numérique** $F_h(u_h) = 0$
3. résolution numérique trop coûteuse : **réduction de modèle** $F_n(u_n) = 0$

Les méthodes de réduction de modèle

Approches qui utilisent des représentations sous **format séparé de faible rang** pour alléger les coûts de calculs lors de la **résolution numérique** d'équations aux dérivées partielles.

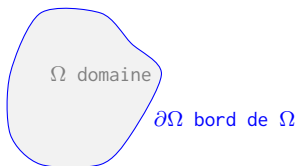
1. Introduction
2. Equation d'advection-diffusion dépendant de paramètres
3. Equation de la chaleur
4. Un peu de théorie
5. Conclusion



Problème modèle (simple).

On cherche $u : \Omega \rightarrow \mathbb{R}$ solution d'une EDP (équation aux dérivées partielles)

$$\left\{ \begin{array}{l} \mathcal{L}u = f, \text{ dans } \Omega, \\ + \text{ conditions aux limites sur } \partial\Omega. \end{array} \right. \quad (1)$$



Problème modèle (simple).

On cherche $u : \Omega \rightarrow \mathbb{R}$ solution d'une EDP (équation aux dérivées partielles)

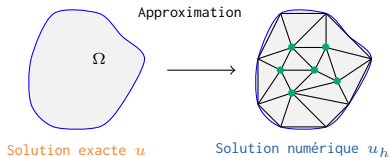
$$\left\{ \begin{array}{l} \mathcal{L}u = f, \text{ dans } \Omega, \\ + \text{ conditions aux limites sur } \partial\Omega. \end{array} \right. \quad (1)$$

avec

- $f : \Omega \rightarrow \mathbb{R}$ est un terme source connu
- $\Omega \subset \mathbb{R}^d$ le domaine spatial
- $\mathcal{L}$ un opérateur différentiel linéaire

ex. : Equation d'advection diffusion $\mathcal{L}u = -\kappa\Delta u + a \cdot \nabla u$

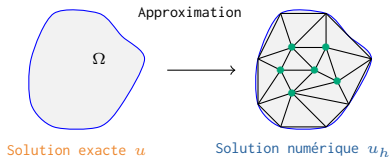
Approximation.



On veut calculer numériquement une approximation u_h de u .

ex. : Méthode éléments finis, différences finies etc.

Approximation.



On veut calculer numériquement une approximation u_h de u .

ex. : Méthode éléments finis, différences finies etc.

Système algébrique de taille N .

En pratique, pour obtenir $u_h \approx u$, il faut calculer le vecteur des coefficients

$$\mathbf{u}_h = (u_1, \dots, u_N)^T \in \mathbb{R}^N$$

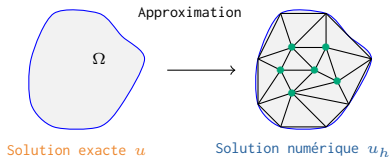
solution de

$$L_h \mathbf{u}_h = \mathbf{f}_h, \text{ dans } \mathbb{R}^N \quad (2)$$

avec $L_h \in \mathbb{R}^{N \times N}$ et $\mathbf{f}_h \in \mathbb{R}^N$ opérateur et terme source discrets.

ex. : Pour les éléments finis $\mathcal{L} = -\frac{d^2}{dx^2} \Rightarrow L_h = \frac{1}{h} \text{tridiag}(-1, 2, -1)$

Approximation.



On veut calculer numériquement une approximation u_h de u .

ex. : Méthode éléments finis, différences finies etc.

Système algébrique de taille N .

En pratique, pour obtenir $u_h \approx u$, il faut calculer le vecteur des coefficients

$$\mathbf{u}_h = (u_1, \dots, u_N)^T \in \mathbb{R}^N$$

solution de

$$L_h \mathbf{u}_h = \mathbf{f}_h, \text{ dans } \mathbb{R}^N \quad (2)$$

avec $L_h \in \mathbb{R}^{N \times N}$ et $\mathbf{f}_h \in \mathbb{R}^N$ opérateur et terme source discrets.

ex. : Pour les éléments finis $\mathcal{L} = -\frac{d^2}{dx^2} \Rightarrow L_h = \frac{1}{h} \text{tridiag}(-1, 2, -1)$

Intérêt ? Ce système linéaire peut être résolu efficacement sur ordinateur.

Combien ça coûte ?

Pour mesurer le coût de calcul on parle de **complexité algorithmique**.

- C'est le nombre d'opérations (+, −, ×, /) pour résoudre ce système linéaire.
- Elle dépend de N !

Combien ça coûte ?

Pour mesurer le coût de calcul on parle de **complexité algorithmique**.

- C'est le nombre d'opérations (+, -, ×, /) pour résoudre ce système linéaire.
- Elle dépend de N !

Quelques algorithmes

- ✗ Méthode de Cramer $\mathcal{O}((N+1)!)$ Si $N = 50$ il faut $51! \approx 1.5 \cdot 10^{66}$ opérations.
 $\approx 5 \cdot 10^{49}$ années sur une machine équipée d'un processeur à $1G.s^{-1}$
- ✓ Méthode du pivot de Gauss $\mathcal{O}(N^3)$
- ✓ Méthode Jacobi, Gauss-Seidel $\mathcal{O}(K \cdot N^2)$

Combien ça coûte ?

Pour mesurer le coût de calcul on parle de **complexité algorithmique**.

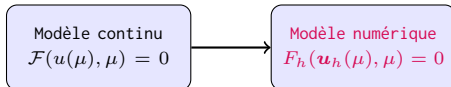
- C'est le nombre d'opérations (+, -, ×, /) pour résoudre ce système linéaire.
- Elle dépend de N !

Quelques algorithmes

- ✗ Méthode de Cramer $\mathcal{O}((N+1)!)$ Si $N = 50$ il faut $50! \approx 1.5 \cdot 10^{66}$ opérations.
 $\approx 5 \cdot 10^{49}$ années sur une machine équipée d'un processeur à $1G.s^{-1}$
- ✓ Méthode du pivot de Gauss $\mathcal{O}(N^3)$
- ✓ Méthode Jacobi, Gauss-Seidel $\mathcal{O}(K \cdot N^2)$

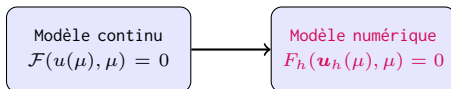
⇒ Pour des **modèles numériques complexes** $N \gg 1$, des algorithmes (solveurs) performants et parallélisés sur des machines très puissantes sont utilisés.

Problème paramétré. On cherche $u : \mu \mapsto u(\mu)$ avec $\mu \in \mathcal{P} \subset \mathbb{R}^p$ des paramètres.



ex. : $-\kappa(\mu)\Delta u(\mu) + a(\mu) \cdot \nabla u(\mu) = f(\mu)$

Problème paramétré. On cherche $u : \mu \mapsto u(\mu)$ avec $\mu \in \mathcal{P} \subset \mathbb{R}^p$ des paramètres.

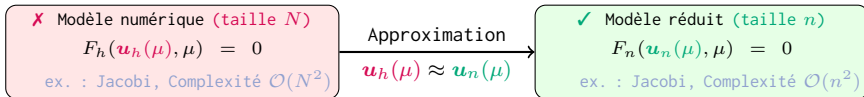


ex. : $-\kappa(\mu)\Delta u(\mu) + a(\mu) \cdot \nabla u(\mu) = f(\mu)$

Problèmes directs.

- Calculer efficacement $\mathbf{u}_h(\mu)$ pour de nombreuses valeurs de $\mu \in \mathcal{P}$.
- **Modèle numérique** $N \gg 1 \rightsquigarrow$ Trop coûteux.

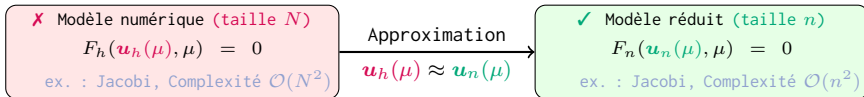
Peut-on remplacer le modèle numérique (N) par autre modèle
dont le coût de résolution est réduit (n) ?



1. Berkooz, G; Holmes, P; Lumley, J L (January 1993). "The Proper Orthogonal Decomposition in the Analysis of Turbulent Flows". Annual Review of Fluid Mechanics. 25 (1): 539-575.

2. A. K. Noor and J. M. Peters. Reduced basis technique for nonlinear analysis of structures. AIAA Journal, 18(4) :455-462, April 1980.

Peut-on remplacer le modèle numérique (N) par autre modèle dont le coût de résolution est réduit (n) ?



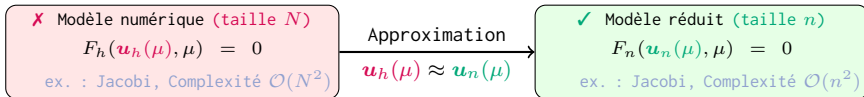
Méthodes de réduction de modèle.

1. **Offline** : Effectuer des (pré)calculs pour construire le modèle réduit.
2. **Online** : On résout le modèle réduit pour calculer des solutions.

1. Berkooz, G; Holmes, P; Lumley, J L (January 1993). "The Proper Orthogonal Decomposition in the Analysis of Turbulent Flows". Annual Review of Fluid Mechanics. 25 (1): 539-575.

2. A. K. Noor and J. M. Peters. Reduced basis technique for nonlinear analysis of structures. AIAA Journal, 18(4) :455-462, April 1980.

Peut-on remplacer le modèle numérique (N) par autre modèle dont le coût de résolution est réduit (n) ?



Méthodes de réduction de modèle.

1. **Offline** : Effectuer des (pré)calculs pour construire le modèle réduit.
2. **Online** : On résoud le modèle réduit pour calculer des solutions.

⇒ Quelques méthodes de réduction de modèle connues :

- Proper Orthogonal Decomposition¹ (POD)
- Bases réduites² (BR)

1. Berkooz, G; Holmes, P; Lumley, J L (January 1993). "The Proper Orthogonal Decomposition in the Analysis of Turbulent Flows". Annual Review of Fluid Mechanics. 25 (1): 539-575.

2. A. K. Noor and J. M. Peters. Reduced basis technique for nonlinear analysis of structures. AIAA Journal, 18(4) :455-462, April 1980.

Les méthodes de réduction de modèle classiques (POD, BR) sont des **méthodes d'approximation linéaire**.

Les méthodes de réduction de modèle classiques (POD, BR) sont des **méthodes d'approximation linéaire**.

On cherche une approximation $\mathbf{u}_n(\mu) \in V_n$ de $\mathbf{u}_h(\mu) \in \mathbb{R}^N$

$$\mathbf{u}_n(\mu) = \sum_{i=1}^n \alpha_i(\mu) \varphi_i \in V_n$$

avec $V_n = \text{vect}\{\varphi_1, \dots, \varphi_n\} \subset \mathbb{R}^N$ et $\dim V_n = n$.

Les méthodes de réduction de modèle classiques (POD, BR) sont des **méthodes d'approximation linéaire**.

On cherche une approximation $\mathbf{u}_n(\mu) \in V_n$ de $\mathbf{u}_h(\mu) \in \mathbb{R}^N$

$$\mathbf{u}_n(\mu) = \sum_{i=1}^n \alpha_i(\mu) \varphi_i \in V_n$$

avec $V_n = \text{vect}\{\varphi_1, \dots, \varphi_n\} \subset \mathbb{R}^N$ et $\dim V_n = n$.

La fonction $\mathbf{u}_n$ est une **approximation de faible rang** de $\mathbf{u}_h$.

① Etape offline : espace réduit

✗ Cette étape peut être coûteuse mais n'est réalisée qu'une seule fois !

① Etape offline : espace réduit

✗ Cette étape peut être coûteuse mais n'est réalisée qu'une seule fois !

Objectif. On veut construire un espace réduit

$$V_n = \text{vect}\{\varphi_1, \dots, \varphi_n\}.$$

✗ Cette étape peut être coûteuse mais n'est réalisée qu'une seule fois !

Objectif. On veut construire un espace réduit

$$V_n = \text{vect}\{\varphi_1, \dots, \varphi_n\}.$$

Pour construire V_n on effectue des "pré-calculs" en résolvant M fois

$$F_h(\mathbf{u}_h(\mu_i), \mu_i) = 0, \quad i = 1, \dots, M,$$

pour M valeurs du paramètre $\{\mu_1, \dots, \mu_M\}$.

On forme la **matrice des snapshots**

$$A = [\mathbf{u}_h(\mu_1), \dots, \mathbf{u}_h(\mu_M)] \in \mathbb{R}^{N \times M}.$$

✗ Cette étape peut être coûteuse mais n'est réalisée qu'une seule fois !

Objectif. On veut construire un espace réduit

$$V_n = \text{vect}\{\varphi_1, \dots, \varphi_n\}.$$

Pour construire V_n on effectue des "pré-calculs" en résolvant M fois

$$F_h(\mathbf{u}_h(\mu_i), \mu_i) = 0, \quad i = 1, \dots, M,$$

pour M valeurs du paramètre $\{\mu_1, \dots, \mu_M\}$.

On forme la **matrice des snapshots**

$$A = [\mathbf{u}_h(\mu_1), \dots, \mathbf{u}_h(\mu_M)] \in \mathbb{R}^{N \times M}.$$

On calcule la **décomposition en valeurs singulières** tronquée à n termes de A .

Wait a minute !

✗ Cette étape peut être coûteuse mais n'est réalisée qu'une seule fois !

Objectif. On veut construire un espace réduit

$$V_n = \text{vect}\{\varphi_1, \dots, \varphi_n\}.$$

Pour construire V_n on effectue des "pré-calculs" en résolvant M fois

$$F_h(\mathbf{u}_h(\mu_i), \mu_i) = 0, \quad i = 1, \dots, M,$$

pour M valeurs du paramètre $\{\mu_1, \dots, \mu_M\}$.

On forme la **matrice des snapshots**

$$A = [\mathbf{u}_h(\mu_1), \dots, \mathbf{u}_h(\mu_M)] \in \mathbb{R}^{N \times M}.$$

On calcule la **décomposition en valeurs singulières** tronquée à n termes de A .

Wait a minute !

$\Rightarrow V_n$ est l'espace engendré par les n premiers **vecteurs singuliers** à gauche de A .

C'est quoi la décomposition en valeurs singulières ?

On veut représenter $A \in \mathbb{R}^{N \times M}$ sous la forme d'un produit de trois matrices

$$A = \Phi \Sigma \Psi^T = \sum_{i=1}^r \sigma_i \varphi_i \psi_i^T.$$

avec

- les vect. singuliers (orthonormaux) : à gauche $\varphi_i \in \mathbb{R}^N$, à droite $\psi_i \in \mathbb{R}^M$,
- les valeurs singulières $\sigma_i \in \mathbb{R}_+$,
- le rang $r = \text{rang}(A) \leq \min(M, N)$.

C'est quoi la décomposition en valeurs singulières ?

On veut représenter $A \in \mathbb{R}^{N \times M}$ sous la forme d'un produit de trois matrices

$$A = \Phi \Sigma \Psi^T = \sum_{i=1}^r \sigma_i \varphi_i \psi_i^T.$$

avec

- les vect. singuliers (orthonormaux) : à gauche $\varphi_i \in \mathbb{R}^N$, à droite $\psi_i \in \mathbb{R}^M$,
- les valeurs singulières $\sigma_i \in \mathbb{R}_+$,
- le rang $r = \text{rang}(A) \leq \min(M, N)$.

Exemple. Pour $N = 6, M = 5$.

$$\underbrace{\begin{pmatrix} a_{11} & a_{12} & a_{13} & a_{14} & a_{15} \\ a_{21} & a_{22} & a_{23} & a_{24} & a_{25} \\ a_{31} & a_{32} & a_{33} & a_{34} & a_{35} \\ a_{41} & a_{42} & a_{43} & a_{44} & a_{45} \\ a_{51} & a_{52} & a_{53} & a_{54} & a_{55} \\ a_{61} & a_{62} & a_{63} & a_{64} & a_{65} \end{pmatrix}}_{A \in \mathbb{R}^{6 \times 5}} = \underbrace{\begin{pmatrix} \varphi_1 & \varphi_2 & \varphi_3 & \varphi_4 & \varphi_5 & \varphi_6 \end{pmatrix}}_{\Phi \in \mathbb{R}^{6 \times 6}} \underbrace{\begin{pmatrix} \sigma_{11} & 0 & 0 & 0 & 0 \\ 0 & \sigma_{22} & 0 & 0 & 0 \\ 0 & 0 & \sigma_{33} & 0 & 0 \\ 0 & 0 & 0 & \sigma_{44} & 0 \\ 0 & 0 & 0 & 0 & \sigma_{55} \\ 0 & 0 & 0 & 0 & 0 \end{pmatrix}}_{\Sigma \in \mathbb{R}^{6 \times 5}} \underbrace{\begin{pmatrix} \psi_1 \\ \psi_2 \\ \psi_3 \\ \psi_4 \\ \psi_5 \end{pmatrix}}_{\Psi^T \in \mathbb{R}^{5 \times 5}}$$

C'est quoi la décomposition en valeurs singulières ?

On veut représenter $A \in \mathbb{R}^{N \times M}$ sous la forme d'un produit de trois matrices

$$A = \Phi \Sigma \Psi^T = \sum_{i=1}^r \sigma_i \varphi_i \psi_i^T.$$

avec

- les vect. singuliers (orthonormaux) : à gauche $\varphi_i \in \mathbb{R}^N$, à droite $\psi_i \in \mathbb{R}^M$,
- les valeurs singulières $\sigma_i \in \mathbb{R}_+$,
- le rang $r = \text{rang}(A) \leq \min(M, N)$.

Exemple. Pour $N = 6, M = 5$.

$$\underbrace{\begin{pmatrix} a_{11} & a_{12} & a_{13} & a_{14} & a_{15} \\ a_{21} & a_{22} & a_{23} & a_{24} & a_{25} \\ a_{31} & a_{32} & a_{33} & a_{34} & a_{35} \\ a_{41} & a_{42} & a_{43} & a_{44} & a_{45} \\ a_{51} & a_{52} & a_{53} & a_{54} & a_{55} \\ a_{61} & a_{62} & a_{63} & a_{64} & a_{65} \end{pmatrix}}_{A \in \mathbb{R}^{6 \times 5}} = \underbrace{\begin{pmatrix} \varphi_1 & \varphi_2 & \varphi_3 & \varphi_4 & \varphi_5 & \varphi_6 \end{pmatrix}}_{\Phi \in \mathbb{R}^{6 \times 6}} \underbrace{\begin{pmatrix} \sigma_{11} & 0 & 0 & 0 & 0 \\ 0 & \sigma_{22} & 0 & 0 & 0 \\ 0 & 0 & \sigma_{33} & 0 & 0 \\ 0 & 0 & 0 & \sigma_{44} & 0 \\ 0 & 0 & 0 & 0 & \sigma_{55} \\ 0 & 0 & 0 & 0 & 0 \end{pmatrix}}_{\Sigma \in \mathbb{R}^{6 \times 5}} \underbrace{\begin{pmatrix} \psi_1 \\ \psi_2 \\ \psi_3 \\ \psi_4 \\ \psi_5 \end{pmatrix}}_{\Psi^T \in \mathbb{R}^{5 \times 5}}$$

Les **trois premiers vecteurs singuliers à gauche** forment un sous-espace réduit de $\mathbb{R}^6$ de dimension $n = 3$:

$$V_n := \text{vect}\{\varphi_1, \varphi_2, \varphi_3\}.$$

Objectif. On veut construire le modèle réduit

$$F_n(\mathbf{u}_n(\mu), \mu) = 0.$$

Objectif. On veut construire le modèle réduit

$$F_n(\mathbf{u}_n(\mu), \mu) = 0.$$

Ansatz. On cherche une approximation $\mathbf{u}_n(\mu)$ de $\mathbf{u}_h(\mu)$ dans V_n i.e.

$$\mathbf{u}_n(\mu) = \sum_{i=1}^n \alpha_i(\mu) \varphi_i = \mathbf{V} \boldsymbol{\alpha}_n(\mu) \in V_n,$$

avec $\mathbf{V} = [\varphi_1, \dots, \varphi_n] \in \mathbb{R}^{N \times n}$ et $\boldsymbol{\alpha}_n(\mu) = (\alpha_1(\mu), \dots, \alpha_n(\mu))^T \in \mathbb{R}^n$.

Objectif. On veut construire le modèle réduit

$$F_n(\mathbf{u}_n(\mu), \mu) = 0.$$

Ansatz. On cherche une approximation $\mathbf{u}_n(\mu)$ de $\mathbf{u}_h(\mu)$ dans V_n i.e.

$$\mathbf{u}_n(\mu) = \sum_{i=1}^n \alpha_i(\mu) \varphi_i = \mathbf{V} \boldsymbol{\alpha}_n(\mu) \in V_n,$$

avec $\mathbf{V} = [\varphi_1, \dots, \varphi_n] \in \mathbb{R}^{N \times n}$ et $\boldsymbol{\alpha}_n(\mu) = (\alpha_1(\mu), \dots, \alpha_n(\mu))^T \in \mathbb{R}^n$.

Projection. Le modèle réduit peut être obtenu par **projection de Galerkin** dans V_n

$$\langle L_h(\mu) \mathbf{u}_n(\mu), \mathbf{v} \rangle_2 = \langle \mathbf{f}_h(\mu), \mathbf{v} \rangle_2, \quad \forall \mathbf{v} \in V_n,$$

Objectif. On veut construire le modèle réduit

$$F_n(\mathbf{u}_n(\mu), \mu) = 0.$$

Ansatz. On cherche une approximation $\mathbf{u}_n(\mu)$ de $\mathbf{u}_h(\mu)$ dans V_n i.e.

$$\mathbf{u}_n(\mu) = \sum_{i=1}^n \alpha_i(\mu) \varphi_i = \mathbf{V} \boldsymbol{\alpha}_n(\mu) \in V_n,$$

avec $\mathbf{V} = [\varphi_1, \dots, \varphi_n] \in \mathbb{R}^{N \times n}$ et $\boldsymbol{\alpha}_n(\mu) = (\alpha_1(\mu), \dots, \alpha_n(\mu))^T \in \mathbb{R}^n$.

Projection. Le modèle réduit peut être obtenu par **projection de Galerkin** dans V_n

$$\langle L_h(\mu) \mathbf{u}_n(\mu), \mathbf{v} \rangle_2 = \langle \mathbf{f}_h(\mu), \mathbf{v} \rangle_2, \quad \forall \mathbf{v} \in V_n,$$

$$\Leftrightarrow \underbrace{\mathbf{V}^T L_h(\mu) \mathbf{V}}_{L_n(\mu) \in \mathbb{R}^{n \times n}} \boldsymbol{\alpha}_n(\mu) = \underbrace{\mathbf{V}^T \mathbf{f}_h(\mu)}_{\mathbf{f}_n(\mu) \in \mathbb{R}^{n \times n}}.$$

On a obtenu un système linéaire de taille $n \times n$!

Système algébrique pour $n \ll N$

Pour n'importe quelle valeur de $\mu \in \mathcal{P}$ on cherche $\alpha_n(\mu) \in \mathbb{R}^n$ t.q.

$$L_n(\mu)\alpha_n(\mu) = f_n(\mu) \text{ dans } \mathbb{R}^n.$$

pour obtenir l'approximation $u_n(\mu) = V\alpha_n(\mu)$ de $u_h(\mu)$.

Système algébrique pour $n \ll N$

Pour n'importe quelle valeur de $\mu \in \mathcal{P}$ on cherche $\alpha_n(\mu) \in \mathbb{R}^n$ t.q.

$$L_n(\mu)\alpha_n(\mu) = f_n(\mu) \text{ dans } \mathbb{R}^n.$$

pour obtenir l'approximation $u_n(\mu) = V\alpha_n(\mu)$ de $u_h(\mu)$.

✓ Cette étape est "cheap" et peut être réalisée autant de fois que possible !

Questions.

1. Quelle est l'erreur commise quand $\mathbf{u}_h(\mu) \approx \mathbf{u}_n(\mu) \in V_n$: $\|\mathbf{u}_h(\mu) - \mathbf{u}_n(\mu)\|_2$?
2. Peut-on la comparer à l'idéal ?

Questions.

1. Quelle est l'erreur commise quand $\mathbf{u}_h(\mu) \approx \mathbf{u}_n(\mu) \in V_n$: $\|\mathbf{u}_h(\mu) - \mathbf{u}_n(\mu)\|_2$?
2. Peut-on la comparer à l'idéal ?

C-à-d la meilleure approximation de $\mathbf{u}_h(\mu)$ pour $\|\cdot\|_2$ dans V_n

$$\|\mathbf{u}_h(\mu) - P_{V_n} \mathbf{u}_h(\mu)\|_2 = \min_{\mathbf{v} \in V_n} \|\mathbf{u}_h(\mu) - \mathbf{v}\|_2.$$

qui est la donnée par la projection orthogonale avec $P_{V_n} = \mathbf{V}\mathbf{V}^T$.

On suppose qu'il existe $0 < \alpha < \beta$ t.q. $L_h(\mu) \in \mathbb{R}^{N \times N}$ vérifie

$$\alpha \|v\|_2^2 \leq \langle L_h(\mu)v, v \rangle_2, \quad \forall v \in \mathbb{R}^N. \quad (3)$$

et

$$\langle L_h(\mu)u, v \rangle_2 \leq \beta \|u\|_2 \|v\|_2, \quad \forall u, v \in \mathbb{R}^N. \quad (4)$$

On suppose qu'il existe $0 < \alpha < \beta$ t.q. $L_h(\mu) \in \mathbb{R}^{N \times N}$ vérifie

$$\alpha \|v\|_2^2 \leq \langle L_h(\mu)v, v \rangle_2, \quad \forall v \in \mathbb{R}^N. \quad (3)$$

et

$$\langle L_h(\mu)u, v \rangle_2 \leq \beta \|u\|_2 \|v\|_2, \quad \forall u, v \in \mathbb{R}^N. \quad (4)$$

Quasi-optimalité

Il existe $\gamma > 0$ (indépendante de n) t.q.

$$\|u_h(\mu) - u_n(\mu)\|_2 = \gamma \min_{v \in V_n} \|u_h(\mu) - v\|_2. \quad (5)$$

Preuve. u_h et u_n satisfont la l'orthogonalité de Galerkin

$$\langle L_h(u_n - u_h), v \rangle_2 = 0, \quad \forall v \in V_n.$$

Preuve. $\mathbf{u}_h$ et $\mathbf{u}_n$ satisfont la l'orthogonalité de Galerkin

$$\langle L_h(\mathbf{u}_n - \mathbf{u}_h), \mathbf{v} \rangle_2 = 0, \quad \forall \mathbf{v} \in V_n.$$

On a

$$\|\mathbf{u}_n - \mathbf{u}_h\|_2^2 \stackrel{(3)}{\leq} \frac{1}{\alpha} \langle L_h(\mathbf{u}_n - \mathbf{u}_h), \mathbf{u}_n - \mathbf{u}_h \rangle_2$$

Preuve. $\mathbf{u}_h$ et $\mathbf{u}_n$ satisfont la l'orthogonalité de Galerkin

$$\langle L_h(\mathbf{u}_n - \mathbf{u}_h), \mathbf{v} \rangle_2 = 0, \quad \forall \mathbf{v} \in V_n.$$

On a

$$\begin{aligned} \|\mathbf{u}_n - \mathbf{u}_h\|_2^2 &\stackrel{(3)}{\leq} \frac{1}{\alpha} \langle L_h(\mathbf{u}_n - \mathbf{u}_h), \mathbf{u}_n - \mathbf{u}_h \rangle_2 \\ \Rightarrow \|\mathbf{u}_n - \mathbf{u}_h\|_2^2 &\leq \frac{1}{\alpha} \langle L_h(\mathbf{u}_n - \mathbf{u}_h), \underbrace{P_{V_n} \mathbf{u}_h - \mathbf{u}_h + \mathbf{u}_n - P_{V_n} \mathbf{u}_h}_{\in V_n} \rangle_2 \end{aligned}$$

Preuve. $\mathbf{u}_h$ et $\mathbf{u}_n$ satisfont la **l'orthogonalité de Galerkin**

$$\langle L_h(\mathbf{u}_n - \mathbf{u}_h), \mathbf{v} \rangle_2 = 0, \quad \forall \mathbf{v} \in V_n.$$

On a

$$\|\mathbf{u}_n - \mathbf{u}_h\|_2^2 \stackrel{(3)}{\leq} \frac{1}{\alpha} \langle L_h(\mathbf{u}_n - \mathbf{u}_h), \mathbf{u}_n - \mathbf{u}_h \rangle_2$$

$$\Rightarrow \|\mathbf{u}_n - \mathbf{u}_h\|_2^2 \leq \frac{1}{\alpha} \langle L_h(\mathbf{u}_n - \mathbf{u}_h), \underbrace{P_{V_n} \mathbf{u}_h - \mathbf{u}_h + \mathbf{u}_n - P_{V_n} \mathbf{u}_h}_{\in V_n} \rangle_2$$

$$\Rightarrow \|\mathbf{u}_n - \mathbf{u}_h\|_2^2 \leq \frac{1}{\alpha} \langle L_h(\mathbf{u}_n - \mathbf{u}_h), P_{V_n} \mathbf{u}_h - \mathbf{u}_h \rangle_2$$

Preuve. $\mathbf{u}_h$ et $\mathbf{u}_n$ satisfont la l'orthogonalité de Galerkin

$$\langle L_h(\mathbf{u}_n - \mathbf{u}_h), \mathbf{v} \rangle_2 = 0, \quad \forall \mathbf{v} \in V_n.$$

On a

$$\|\mathbf{u}_n - \mathbf{u}_h\|_2^2 \stackrel{(3)}{\leq} \frac{1}{\alpha} \langle L_h(\mathbf{u}_n - \mathbf{u}_h), \mathbf{u}_n - \mathbf{u}_h \rangle_2$$

$$\Rightarrow \|\mathbf{u}_n - \mathbf{u}_h\|_2^2 \leq \frac{1}{\alpha} \langle L_h(\mathbf{u}_n - \mathbf{u}_h), \mathbf{P}_{V_n} \mathbf{u}_h - \mathbf{u}_h + \underbrace{\mathbf{u}_n - \mathbf{P}_{V_n} \mathbf{u}_h}_{\in V_n} \rangle_2$$

$$\Rightarrow \|\mathbf{u}_n - \mathbf{u}_h\|_2^2 \leq \frac{1}{\alpha} \langle L_h(\mathbf{u}_n - \mathbf{u}_h), \mathbf{P}_{V_n} \mathbf{u}_h - \mathbf{u}_h \rangle_2$$

$$\Rightarrow \|\mathbf{u}_n - \mathbf{u}_h\|_2^2 \stackrel{(4)}{\leq} \frac{\beta}{\alpha} \|\mathbf{u}_n - \mathbf{u}_h\|_2 \|\mathbf{P}_{V_n} \mathbf{u}_h - \mathbf{u}_h\|_2$$

Preuve. $\mathbf{u}_h$ et $\mathbf{u}_n$ satisfont la l'orthogonalité de Galerkin

$$\langle L_h(\mathbf{u}_n - \mathbf{u}_h), \mathbf{v} \rangle_2 = 0, \quad \forall \mathbf{v} \in V_n.$$

On a

$$\|\mathbf{u}_n - \mathbf{u}_h\|_2^2 \stackrel{(3)}{\leq} \frac{1}{\alpha} \langle L_h(\mathbf{u}_n - \mathbf{u}_h), \mathbf{u}_n - \mathbf{u}_h \rangle_2$$

$$\Rightarrow \|\mathbf{u}_n - \mathbf{u}_h\|_2^2 \leq \frac{1}{\alpha} \langle L_h(\mathbf{u}_n - \mathbf{u}_h), \mathbf{P}_{V_n} \mathbf{u}_h - \mathbf{u}_h + \underbrace{\mathbf{u}_n - \mathbf{P}_{V_n} \mathbf{u}_h}_{\in V_n} \rangle_2$$

$$\Rightarrow \|\mathbf{u}_n - \mathbf{u}_h\|_2^2 \leq \frac{1}{\alpha} \langle L_h(\mathbf{u}_n - \mathbf{u}_h), \mathbf{P}_{V_n} \mathbf{u}_h - \mathbf{u}_h \rangle_2$$

$$\Rightarrow \|\mathbf{u}_n - \mathbf{u}_h\|_2^2 \stackrel{(4)}{\leq} \frac{\beta}{\alpha} \|\mathbf{u}_n - \mathbf{u}_h\|_2 \|\mathbf{P}_{V_n} \mathbf{u}_h - \mathbf{u}_h\|_2$$

D'où le résultat (5) en simplifiant et posant $\gamma = \frac{\beta}{\alpha}$.

Illustration

On cherche $u(\mu) : \mathbb{R} \rightarrow \mathbb{R}$ solution de

$$\begin{cases} -\mu_2 u''(x; \mu) + \mu_1 u'(x; \mu) &= f(x), & x \in (0, 1), \\ u(1; \mu) = u(0; \mu) &= 0. \end{cases} \quad (6)$$

On a $\mu = (\mu_1, \mu_2) \in \mathcal{P}$ avec $\mathcal{P} = [1, 10] \times [0.01, 1]$ et $f(x) = 1$. Ici $N = 1000$.

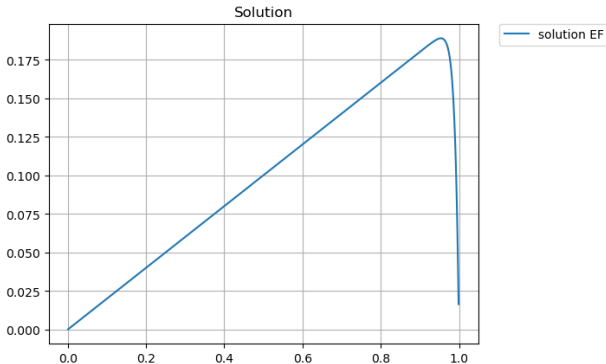


Figure - Solution numérique pour $\mu = (5, 0.05)$

Illustration

On cherche $u(\mu) : \mathbb{R} \rightarrow \mathbb{R}$ solution de

$$\begin{cases} -\mu_2 u''(x; \mu) + \mu_1 u'(x; \mu) &= f(x), & x \in (0, 1), \\ u(1; \mu) = u(0; \mu) &= 0. \end{cases} \quad (6)$$

On a $\mu = (\mu_1, \mu_2) \in \mathcal{P}$ avec $\mathcal{P} = [1, 10] \times [0.01, 1]$ et $f(x) = 1$. Ici $N = 1000$.

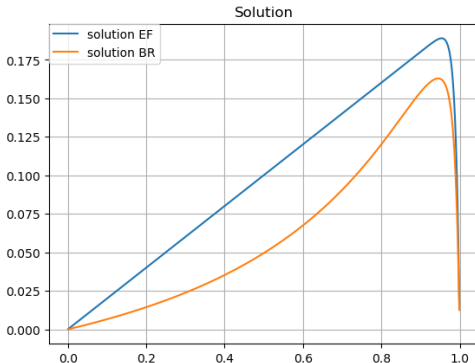


Figure - Base apprise pour $M = 50$ valeurs de μ . Solution calculée pour $n = 2$ et $\mu = (5, 0.05)$

Erreur relative pour $n = 2$: $err_2(\mu) = 0.6663235211802049$

Erreur relative pour $n = 20$: $err_{20}(\mu) = 3.037748791261398e - 07$

Illustration

On cherche $u(\mu) : \mathbb{R} \rightarrow \mathbb{R}$ solution de

$$\begin{cases} -\mu_2 u''(x; \mu) + \mu_1 u'(x; \mu) &= f(x), & x \in (0, 1), \\ u(1; \mu) = u(0; \mu) &= 0. \end{cases} \quad (6)$$

On a $\mu = (\mu_1, \mu_2) \in \mathcal{P}$ avec $\mathcal{P} = [1, 10] \times [0.01, 1]$ et $f(x) = 1$. Ici $N = 1000$.

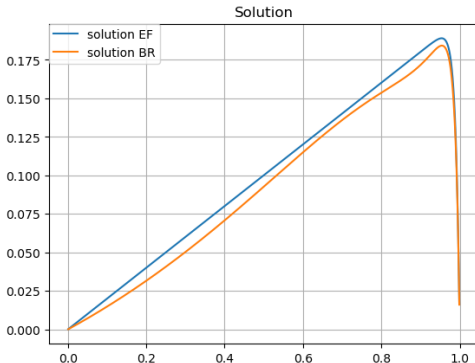


Figure - Base apprise pour $M = 50$ valeurs de μ . Solution calculée pour $n = 5$ et $\mu = (5, 0.05)$

Erreur relative pour $n = 2$: $err_2(\mu) = 0.6663235211802049$

Erreur relative pour $n = 20$: $err_{20}(\mu) = 3.037748791261398e - 07$

Illustration

On cherche $u(\mu) : \mathbb{R} \rightarrow \mathbb{R}$ solution de

$$\begin{cases} -\mu_2 u''(x; \mu) + \mu_1 u'(x; \mu) &= f(x), & x \in (0, 1), \\ u(1; \mu) = u(0; \mu) &= 0. \end{cases} \quad (6)$$

On a $\mu = (\mu_1, \mu_2) \in \mathcal{P}$ avec $\mathcal{P} = [1, 10] \times [0.01, 1]$ et $f(x) = 1$. Ici $N = 1000$.

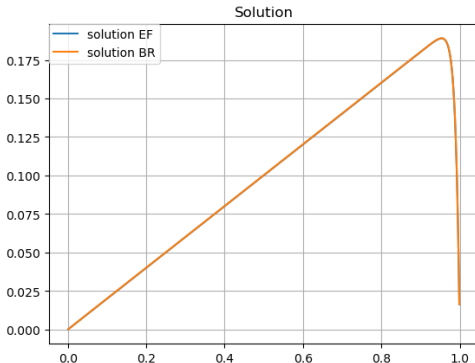


Figure - Base apprise pour $M = 50$ valeurs de μ . Solution calculée pour $n = 10$ et $\mu = (5, 0.05)$

Erreur relative pour $n = 2$: $err_2(\mu) = 0.6663235211802049$

Erreur relative pour $n = 20$: $err_{20}(\mu) = 3.037748791261398e - 07$

Illustration

On cherche $u(\mu) : \mathbb{R} \rightarrow \mathbb{R}$ solution de

$$\begin{cases} -\mu_2 u''(x; \mu) + \mu_1 u'(x; \mu) &= f(x), & x \in (0, 1), \\ u(1; \mu) = u(0; \mu) &= 0. \end{cases} \quad (6)$$

On a $\mu = (\mu_1, \mu_2) \in \mathcal{P}$ avec $\mathcal{P} = [1, 10] \times [0.01, 1]$ et $f(x) = 1$. Ici $N = 1000$.

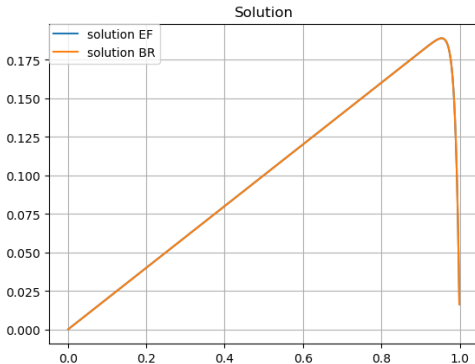


Figure - Base apprise pour $M = 50$ valeurs de μ . Solution calculée pour $n = 20$ et $\mu = (5, 0.05)$

Erreur relative pour $n = 2$: $err_2(\mu) = 0.6663235211802049$

Erreur relative pour $n = 20$: $err_{20}(\mu) = 3.037748791261398e - 07$

1. Introduction
2. Equation d'advection-diffusion dépendant de paramètres
3. Equation de la chaleur
4. Un peu de théorie
5. Conclusion

Problème modèle (simple).

On cherche $u : [0, T] \times \Omega \rightarrow \mathbb{R}$ solution d'une EDP d'évolution

$$\left\{ \begin{array}{l} \partial_t u = \mathcal{L}u, \text{ dans } (0, T] \times \Omega, \\ u(0, \cdot) = u_0, \text{ dans } \Omega, \\ + \text{ conditions aux limites sur } \partial\Omega. \end{array} \right. \quad (7)$$

Problème modèle (simple).

On cherche $u : [0, T] \times \Omega \rightarrow \mathbb{R}$ solution d'une EDP d'évolution

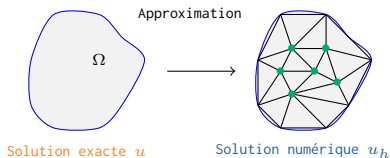
$$\left\{ \begin{array}{l} \partial_t u = \mathcal{L}u, \text{ dans } (0, T] \times \Omega, \\ u(0, \cdot) = u_0, \text{ dans } \Omega, \\ + \text{ conditions aux limites sur } \partial\Omega. \end{array} \right. \quad (7)$$

avec

- $\Omega \subset \mathbb{R}^d$ le domaine spatial
- $\mathcal{L}$ un opérateur différentiel linéaire

ex. : Equation de la chaleur $\mathcal{L}u = \nu \Delta u$

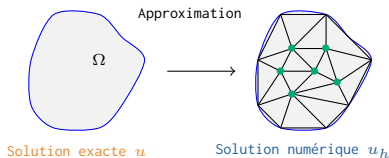
Approximation.



On veut calculer numériquement une approximation $u_h(t)$ de $u(t)$.

ex. : Méthode éléments finis, différences finies etc.

Approximation.



On veut calculer numériquement une approximation $u_h(t)$ de $u(t)$.

ex. : Méthode éléments finis, différences finies etc.

Système d'EDO de taille N .

En pratique, pour obtenir $u_h \approx u$, il faut calculer le vecteur des coefficients

$$\mathbf{u}_h : [0, T] \rightarrow \mathbb{R}^N$$

solution de

$$\frac{d}{dt} \mathbf{u}_h(t) = L_h \mathbf{u}_h(t), \text{ dans } \mathbb{R}^N \quad (8)$$

avec $L_h \in \mathbb{R}^{N \times N}$.

ex. : Pour les différences finies : $\mathcal{L} = \frac{\partial^2}{\partial x^2} \Rightarrow L_h = -\frac{1}{h^2} \text{tridiag}(-1, 2, -1)$.

L'espace réduit

$$V_n = \text{vect}\{\varphi_1, \dots, \varphi_n\}$$

est engendré par les n premiers **vecteurs singuliers** à gauche de

$$A = [\mathbf{u}_h(t_1), \dots, \mathbf{u}_h(t_M)] \in \mathbb{R}^{N \times M}.$$

pour M instants $\{t_1, \dots, t_M\} \subset [0, T]$.

L'espace réduit

$$V_n = \text{vect}\{\varphi_1, \dots, \varphi_n\}$$

est engendré par les n premiers **vecteurs singuliers** à gauche de

$$A = [\mathbf{u}_h(t_1), \dots, \mathbf{u}_h(t_M)] \in \mathbb{R}^{N \times M}.$$

pour M instants $\{t_1, \dots, t_M\} \subset [0, T]$.

On cherche une approximation $\mathbf{u}_n(t) = \mathbf{V}\boldsymbol{\alpha}_n(t) \in V_n$ solution de

$$F_n(u_n) = 0,$$

L'espace réduit

$$V_n = \text{vect}\{\varphi_1, \dots, \varphi_n\}$$

est engendré par les n premiers **vecteurs singuliers** à gauche de

$$A = [\mathbf{u}_h(t_1), \dots, \mathbf{u}_h(t_M)] \in \mathbb{R}^{N \times M}.$$

pour M instants $\{t_1, \dots, t_M\} \subset [0, T]$.

On cherche une approximation $\mathbf{u}_n(t) = \mathbf{V}\boldsymbol{\alpha}_n(t) \in V_n$ solution de

$$F_n(u_n) = 0,$$

obtenu par **projection de Galerkin** dans V_n

$$\frac{d}{dt}\boldsymbol{\alpha}_n(t) = \underbrace{\mathbf{V}^T L_h \mathbf{V}}_{L_n \in \mathbb{R}^{n \times n}} \boldsymbol{\alpha}_n(t).$$

On a obtenu un système d'EDOs de taille $n \times n$!

Système d'EDO de taille $n \ll N$

On cherche $\alpha_n : [0, T] \rightarrow \mathbb{R}^n$ solution de

$$\frac{d}{dt} \alpha_n(t) = L_n \alpha_n(t), t \in (0, T]$$

avec la condition initiale $\alpha(0) = V^T u_{h,0}$.

Soit ℓ la constante de lipschitz logarithmique pour $L_h(\mu) \in \mathbb{R}^{N \times N}$ t.q.

$$\ell := \sup_{\mathbf{v} \neq 0} \frac{\langle L_h \mathbf{v}, \mathbf{v} \rangle_2}{\|\mathbf{v}\|_2} = \lambda_{\max} \left(\frac{1}{2} (L_h + L_h^T) \right). \quad (9)$$

Borne de l'erreur

Pour tout $t \in [0, T]$, on a

$$\|\mathbf{u}_h(t) - \mathbf{u}_n(t)\|_2 \leq e^{\ell t} \underbrace{\|\mathbf{u}_n(0) - \mathbf{u}_h(0)\|_2}_{\textcircled{1}} + \int_0^t e^{\ell(t-s)} \underbrace{\|(\mathbf{I}_N - \mathbf{P}_{V_n}) L_h \mathbf{u}_n(s)\|_2}_{\textcircled{2}} ds, \quad (10)$$

avec les erreurs de projection dans V_n

- $\textcircled{1}$ de la condition initiale : $\mathbf{u}_n(0) = \mathbf{P}_{V_n} \mathbf{u}_h(0)$,
- $\textcircled{2}$ du "second membre" : $(\mathbf{I}_N - \mathbf{P}_{V_n}) L_h \mathbf{u}_n$.

Preuve. Rappelons que pour tout $t \in (0, T]$ on a

$$\frac{d}{dt} \mathbf{u}_h = L_h \mathbf{u}_h, \quad \frac{d}{dt} \mathbf{u}_n = P_{V_n} L_h \mathbf{u}_n.$$

Preuve. Rappelons que pour tout $t \in (0, T]$ on a

$$\frac{d}{dt}\mathbf{u}_h = L_h\mathbf{u}_h, \quad \frac{d}{dt}\mathbf{u}_n = P_{V_n}L_h\mathbf{u}_n.$$

En soustrayant membre à membre

$$\frac{d}{dt}(\mathbf{u}_h - \mathbf{u}_n) = L_h\mathbf{u}_h - P_{V_n}L_h\mathbf{u}_n$$

Preuve. Rappelons que pour tout $t \in (0, T]$ on a

$$\frac{d}{dt} \mathbf{u}_h = L_h \mathbf{u}_h, \quad \frac{d}{dt} \mathbf{u}_n = P_{V_n} L_h \mathbf{u}_n.$$

En soustrayant membre à membre

$$\begin{aligned} \frac{d}{dt} (\mathbf{u}_h - \mathbf{u}_n) &= L_h \mathbf{u}_h - P_{V_n} L_h \mathbf{u}_n \\ \Rightarrow \frac{d}{dt} (\mathbf{u}_h - \mathbf{u}_n) &= L_h (\mathbf{u}_h - \mathbf{u}_n) + (\mathbf{I}_N - P_{V_n}) L_h \mathbf{u}_n \end{aligned}$$

Preuve. Rappelons que pour tout $t \in (0, T]$ on a

$$\frac{d}{dt} \mathbf{u}_h = L_h \mathbf{u}_h, \quad \frac{d}{dt} \mathbf{u}_n = P_{V_n} L_h \mathbf{u}_n.$$

En soustrayant membre à membre

$$\frac{d}{dt} (\mathbf{u}_h - \mathbf{u}_n) = L_h \mathbf{u}_h - P_{V_n} L_h \mathbf{u}_n$$

$$\Rightarrow \frac{d}{dt} (\mathbf{u}_h - \mathbf{u}_n) = L_h (\mathbf{u}_h - \mathbf{u}_n) + (I_N - P_{V_n}) L_h \mathbf{u}_n$$

$$\Rightarrow \left\langle \frac{d}{dt} (\mathbf{u}_h - \mathbf{u}_n), \mathbf{u}_h - \mathbf{u}_n \right\rangle_2 = \langle L_h (\mathbf{u}_h - \mathbf{u}_n), \mathbf{u}_h - \mathbf{u}_n \rangle_2 + \langle (I_N - P_{V_n}) L_h \mathbf{u}_n, \mathbf{u}_h - \mathbf{u}_n \rangle_2.$$

Preuve. Rappelons que pour tout $t \in (0, T]$ on a

$$\frac{d}{dt} \mathbf{u}_h = L_h \mathbf{u}_h, \quad \frac{d}{dt} \mathbf{u}_n = P_{V_n} L_h \mathbf{u}_n.$$

En soustrayant membre à membre

$$\frac{d}{dt} (\mathbf{u}_h - \mathbf{u}_n) = L_h \mathbf{u}_h - P_{V_n} L_h \mathbf{u}_n$$

$$\Rightarrow \frac{d}{dt} (\mathbf{u}_h - \mathbf{u}_n) = L_h (\mathbf{u}_h - \mathbf{u}_n) + (I_N - P_{V_n}) L_h \mathbf{u}_n$$

$$\Rightarrow \left\langle \frac{d}{dt} (\mathbf{u}_h - \mathbf{u}_n), \mathbf{u}_h - \mathbf{u}_n \right\rangle_2 = \langle L_h (\mathbf{u}_h - \mathbf{u}_n), \mathbf{u}_h - \mathbf{u}_n \rangle_2 + \langle (I_N - P_{V_n}) L_h \mathbf{u}_n, \mathbf{u}_h - \mathbf{u}_n \rangle_2.$$

$$\Rightarrow \|\mathbf{u}_h - \mathbf{u}_n\|_2 \frac{d}{dt} \|\mathbf{u}_h - \mathbf{u}_n\|_2 = \langle L_h (\mathbf{u}_h - \mathbf{u}_n), \mathbf{u}_h - \mathbf{u}_n \rangle_2 - \langle (I_N - P_{V_n}) L_h \mathbf{u}_n, \mathbf{u}_h - \mathbf{u}_n \rangle_2.$$

Estimation d'erreur a priori $\mathbf{u}_h(t) \approx \mathbf{u}_n(t)$

Preuve. Rappelons que pour tout $t \in (0, T]$ on a

$$\frac{d}{dt} \mathbf{u}_h = L_h \mathbf{u}_h, \quad \frac{d}{dt} \mathbf{u}_n = P_{V_n} L_h \mathbf{u}_n.$$

En soustrayant membre à membre

$$\frac{d}{dt} (\mathbf{u}_h - \mathbf{u}_n) = L_h \mathbf{u}_h - P_{V_n} L_h \mathbf{u}_n$$

$$\Rightarrow \frac{d}{dt} (\mathbf{u}_h - \mathbf{u}_n) = L_h (\mathbf{u}_h - \mathbf{u}_n) + (I_N - P_{V_n}) L_h \mathbf{u}_n$$

$$\Rightarrow \left\langle \frac{d}{dt} (\mathbf{u}_h - \mathbf{u}_n), \mathbf{u}_h - \mathbf{u}_n \right\rangle_2 = \langle L_h (\mathbf{u}_h - \mathbf{u}_n), \mathbf{u}_h - \mathbf{u}_n \rangle_2 + \langle (I_N - P_{V_n}) L_h \mathbf{u}_n, \mathbf{u}_h - \mathbf{u}_n \rangle_2.$$

$$\Rightarrow \|\mathbf{u}_h - \mathbf{u}_n\|_2 \frac{d}{dt} \|\mathbf{u}_h - \mathbf{u}_n\|_2 = \langle L_h (\mathbf{u}_h - \mathbf{u}_n), \mathbf{u}_h - \mathbf{u}_n \rangle_2 - \langle (I_N - P_{V_n}) L_h \mathbf{u}_n, \mathbf{u}_h - \mathbf{u}_n \rangle_2.$$

$$\Rightarrow \|\mathbf{u}_h - \mathbf{u}_n\|_2 \frac{d}{dt} \|\mathbf{u}_h - \mathbf{u}_n\|_2 \stackrel{(9)}{\leq} \ell \|\mathbf{u}_h - \mathbf{u}_n\|_2^2 + \|(I_N - P_{V_n}) L_h \mathbf{u}_n\| \|\mathbf{u}_h - \mathbf{u}_n\|_2.$$

Estimation d'erreur a priori $\mathbf{u}_h(t) \approx \mathbf{u}_n(t)$

Preuve. Rappelons que pour tout $t \in (0, T]$ on a

$$\frac{d}{dt} \mathbf{u}_h = L_h \mathbf{u}_h, \quad \frac{d}{dt} \mathbf{u}_n = P_{V_n} L_h \mathbf{u}_n.$$

En soustrayant membre à membre

$$\frac{d}{dt} (\mathbf{u}_h - \mathbf{u}_n) = L_h \mathbf{u}_h - P_{V_n} L_h \mathbf{u}_n$$

$$\Rightarrow \frac{d}{dt} (\mathbf{u}_h - \mathbf{u}_n) = L_h (\mathbf{u}_h - \mathbf{u}_n) + (I_N - P_{V_n}) L_h \mathbf{u}_n$$

$$\Rightarrow \left\langle \frac{d}{dt} (\mathbf{u}_h - \mathbf{u}_n), \mathbf{u}_h - \mathbf{u}_n \right\rangle_2 = \langle L_h (\mathbf{u}_h - \mathbf{u}_n), \mathbf{u}_h - \mathbf{u}_n \rangle_2 + \langle (I_N - P_{V_n}) L_h \mathbf{u}_n, \mathbf{u}_h - \mathbf{u}_n \rangle_2.$$

$$\Rightarrow \|\mathbf{u}_h - \mathbf{u}_n\|_2 \frac{d}{dt} \|\mathbf{u}_h - \mathbf{u}_n\|_2 = \langle L_h (\mathbf{u}_h - \mathbf{u}_n), \mathbf{u}_h - \mathbf{u}_n \rangle_2 - \langle (I_N - P_{V_n}) L_h \mathbf{u}_n, \mathbf{u}_h - \mathbf{u}_n \rangle_2.$$

$$\Rightarrow \|\mathbf{u}_h - \mathbf{u}_n\|_2 \frac{d}{dt} \|\mathbf{u}_h - \mathbf{u}_n\|_2 \stackrel{(9)}{\leq} \ell \|\mathbf{u}_h - \mathbf{u}_n\|_2^2 + \|(I_N - P_{V_n}) L_h \mathbf{u}_n\| \|\mathbf{u}_h - \mathbf{u}_n\|_2.$$

$$\Rightarrow \frac{d}{dt} \|\mathbf{u}_h - \mathbf{u}_n\|_2 \leq \ell \|\mathbf{u}_h - \mathbf{u}_n\|_2 + \|(I_N - P_{V_n}) L_h \mathbf{u}_n\|.$$

On obtient le résultat par le lemme de comparaison.

Soit $\{t^0, \dots, t^K\}$ un maillage uniforme de $[0, T]$: $t^k = k\delta t, \delta t = \frac{T}{K}$.

Soit $\{t^0, \dots, t^K\}$ un maillage uniforme de $[0, T]$: $t^k = k\delta t, \delta t = \frac{T}{K}$.

On cherche une approximation de la fonction $\mathbf{u}_n$ en t^k i.e.

$$\mathbf{u}_n(t^k) \approx \mathbf{u}_n^k := \mathbf{V}\boldsymbol{\alpha}_n^k$$

e.g. par schéma **d'Euler implicite** qui permet de construire $\{\boldsymbol{\alpha}_n^k\}_{k=0}^K$ en résolvant

$$(\mathbf{I}_N - \delta t \mathbf{L}_n) \boldsymbol{\alpha}_n^{k+1} = \boldsymbol{\alpha}_n^k, \quad k = 0, \dots, K-1$$

avec $\boldsymbol{\alpha}^0 = \boldsymbol{\alpha}(0)$.

Soit $\{t^0, \dots, t^K\}$ un maillage uniforme de $[0, T]$: $t^k = k\delta t, \delta t = \frac{T}{K}$.

On cherche une approximation de la fonction $\mathbf{u}_n$ en t^k i.e.

$$\mathbf{u}_n(t^k) \approx \mathbf{u}_n^k := \mathbf{V} \boldsymbol{\alpha}_n^k$$

e.g. par schéma **d'Euler implicite** qui permet de construire $\{\boldsymbol{\alpha}_n^k\}_{k=0}^K$ en résolvant

$$(\mathbf{I}_N - \delta t L_n) \boldsymbol{\alpha}_n^{k+1} = \boldsymbol{\alpha}_n^k, \quad k = 0, \dots, K-1$$

avec $\boldsymbol{\alpha}^0 = \boldsymbol{\alpha}(0)$.

✓ On résout une succession de petits problèmes linéaires de taille $n \times n$, K fois !

On cherche $u : [0, T] \times \Omega \rightarrow \mathbb{R}$ solution de

$$\begin{cases} \partial_t u(t, x) &= \partial_{xx}^2 u(t, x), & x \in \Omega, \\ u(t, 0) &= u(t, 1) = 0, \end{cases} \quad (11)$$

avec $u(0, x) = \mathbf{1}_{[0.1, 0.3]}(x) - \mathbf{1}_{[0.6, 0.8]}(x)$, $T = 0.5$. Ici $K = 200$, $N = 1000$.

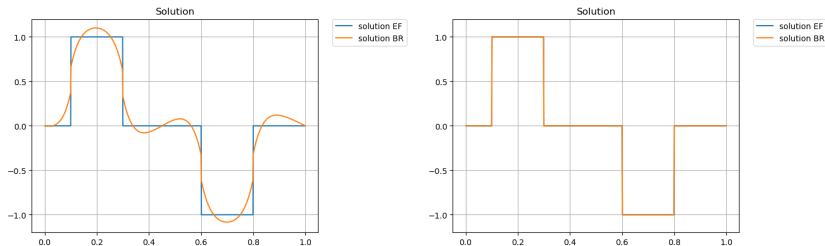


Figure - Base apprise pour $M = 50$. Solution à $t = 0$ pour $n = 2$ (gauche) $n = 20$ (droite).

Erreur relative pour $n = 2$: $err_2 = 0.17475296041661006$
 Erreur relative pour $n = 20$: $err_{20} = 1.7279662844846755e - 05$

On cherche $u : [0, T] \times \Omega \rightarrow \mathbb{R}$ solution de

$$\begin{cases} \partial_t u(t, x) &= \partial_{xx}^2 u(t, x), & x \in \Omega, \\ u(t, 0) &= u(t, 1) = 0, \end{cases} \quad (11)$$

avec $u(0, x) = \mathbf{1}_{[0.1, 0.3]}(x) - \mathbf{1}_{[0.6, 0.8]}(x)$, $T = 0.5$. Ici $K = 200$, $N = 1000$.

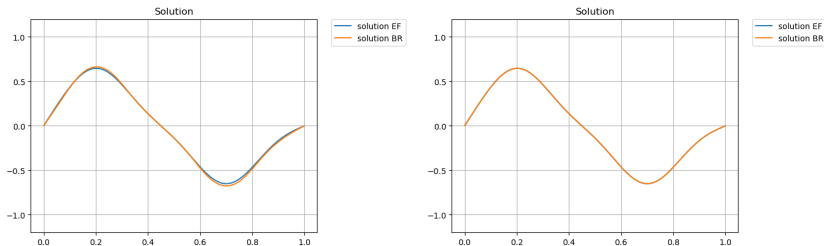


Figure - Base apprise pour $M = 50$. Solution à $t = 0.05$ pour $n = 2$ (gauche) $n = 20$ (droite).

Erreur relative pour $n = 2$: $err_2 = 0.17475296041661006$

Erreur relative pour $n = 20$: $err_{20} = 1.7279662844846755e - 05$

On cherche $u : [0, T] \times \Omega \rightarrow \mathbb{R}$ solution de

$$\begin{cases} \partial_t u(t, x) &= \partial_{xx}^2 u(t, x), & x \in \Omega, \\ u(t, 0) &= u(t, 1) = 0, \end{cases} \quad (11)$$

avec $u(0, x) = \mathbf{1}_{[0.1, 0.3]}(x) - \mathbf{1}_{[0.6, 0.8]}(x)$, $T = 0.5$. Ici $K = 200$, $N = 1000$.

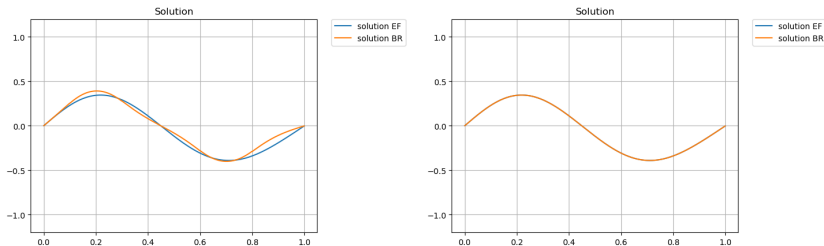


Figure - Base apprise pour $M = 50$. Solution à $t = 0.15$ pour $n = 2$ (gauche) $n = 20$ (droite).

Erreur relative pour $n = 2$: $err_2 = 0.17475296041661006$

Erreur relative pour $n = 20$: $err_{20} = 1.7279662844846755e - 05$

On cherche $u : [0, T] \times \Omega \rightarrow \mathbb{R}$ solution de

$$\begin{cases} \partial_t u(t, x) &= \partial_{xx}^2 u(t, x), & x \in \Omega, \\ u(t, 0) &= u(t, 1) = 0, \end{cases} \quad (11)$$

avec $u(0, x) = \mathbf{1}_{[0.1, 0.3]}(x) - \mathbf{1}_{[0.6, 0.8]}(x)$, $T = 0.5$. Ici $K = 200$, $N = 1000$.

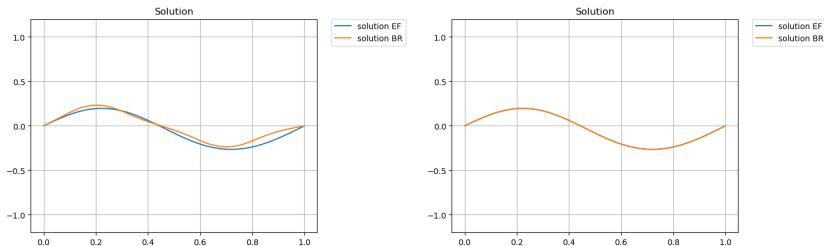


Figure - Base apprise pour $M = 50$. Solution à $t = 0.25$ pour $n = 2$ (gauche) $n = 20$ (droite).

Erreur relative pour $n = 2$: $err_2 = 0.17475296041661006$

Erreur relative pour $n = 20$: $err_{20} = 1.7279662844846755e - 05$

On cherche $u : [0, T] \times \Omega \rightarrow \mathbb{R}$ solution de

$$\begin{cases} \partial_t u(t, x) &= \partial_{xx}^2 u(t, x), & x \in \Omega, \\ u(t, 0) &= u(t, 1) = 0, \end{cases} \quad (11)$$

avec $u(0, x) = \mathbf{1}_{[0.1, 0.3]}(x) - \mathbf{1}_{[0.6, 0.8]}(x)$, $T = 0.5$. Ici $K = 200$, $N = 1000$.

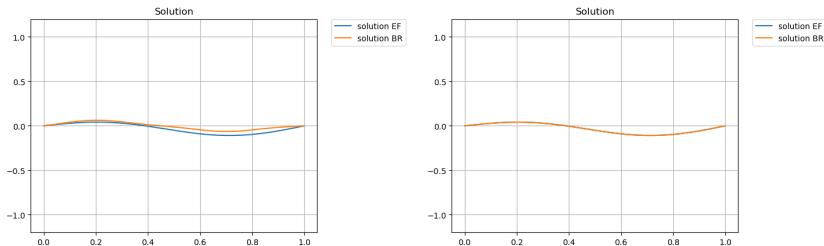


Figure - Base apprise pour $M = 50$. Solution à $t = 0.5$ pour $n = 2$ (gauche) $n = 20$ (droite).

Erreur relative pour $n = 2$: $err_2 = 0.17475296041661006$

Erreur relative pour $n = 20$: $err_{20} = 1.7279662844846755e - 05$

1. Introduction
2. Equation d'advection-diffusion dépendant de paramètres
3. Equation de la chaleur
4. Un peu de théorie
5. Conclusion

Bilan.

Equation de la chaleur

$$\mathbf{u}_n(t) = \sum_{i=1}^n \alpha_i(t) \varphi_i \in V_n$$

Equation d'advection diffusion

$$\mathbf{u}_n(\mu) = \sum_{i=1}^n \alpha_i(\mu) \varphi_i \in V_n$$

Bilan.

Equation de la chaleur

$$\mathbf{u}_n(t) = \sum_{i=1}^n \alpha_i(t) \varphi_i \in V_n$$

Equation d'advection diffusion

$$\mathbf{u}_n(\mu) = \sum_{i=1}^n \alpha_i(\mu) \varphi_i \in V_n$$

Question. Pourquoi une approximation sous format $\mathbf{u}_n \approx \mathbf{u}_h$ de rang n "fonctionne" ?

Bilan.

Equation de la chaleur

$$\mathbf{u}_n(t) = \sum_{i=1}^n \alpha_i(t) \varphi_i \in V_n$$

Equation d'advection diffusion

$$\mathbf{u}_n(\mu) = \sum_{i=1}^n \alpha_i(\mu) \varphi_i \in V_n$$

Question. Pourquoi une approximation sous format $\mathbf{u}_n \approx \mathbf{u}_h$ de rang n "fonctionne" ?

Réponse. On cherche des approximations par le biais de projections dans des espaces V_n (quasi)-optimaux.

Optimalité de la base réduite [QMN, §6.3.1.]

Soit

$$\mathcal{V}_n := \{\mathbf{W} \in \mathbb{R}^{N \times n} : \mathbf{W}\mathbf{W}^T = \mathbf{I}_n\}$$

l'ensemble des matrices orthonormées. Alors

$$\sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - \Phi_n \Phi_n^T \mathbf{u}_h(\mu_i)\|_2^2 = \min_{\mathbf{W} \in \mathcal{V}_n} \sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - \mathbf{W}\mathbf{W}^T \mathbf{u}_h(\mu_i)\|_2^2 = \sum_{i=n+1}^r \sigma_i^2, \quad (12)$$

Optimalité de la base réduite [QMN, §6.3.1.]

Soit

$$\mathcal{V}_n := \{\mathbf{W} \in \mathbb{R}^{N \times n} : \mathbf{W}\mathbf{W}^T = \mathbf{I}_n\}$$

l'ensemble des matrices orthonormées. Alors

$$\sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - \Phi_n \Phi_n^T \mathbf{u}_h(\mu_i)\|_2^2 = \min_{\mathbf{W} \in \mathcal{V}_n} \sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - \mathbf{W}\mathbf{W}^T \mathbf{u}_h(\mu_i)\|_2^2 = \sum_{i=n+1}^r \sigma_i^2, \quad (12)$$

avec

- $\Phi_n = [\varphi_1, \dots, \varphi_n] \in \mathbb{R}^{N \times n}$,
- $\{\sigma_i\}_{i=1}^r \subset \mathbb{R}_+$ les valeurs singulières de
- la matrice de snapshots $A = [\mathbf{u}_h(\mu_1), \dots, \mathbf{u}_h(\mu_M)] \in \mathbb{R}^{N \times n}$.

Interprétation. Pour tout $\mathbf{W} \in \mathcal{V}_n$, on a $P_V = \mathbf{W}\mathbf{W}^T$ pour $V = \text{vectcol}(\mathbf{W})$.

Interprétation. Pour tout $\mathbf{W} \in \mathcal{V}_n$, on a $P_V = \mathbf{W}\mathbf{W}^T$ pour $V = \text{vectcol}(\mathbf{W})$.

Considérons

$$\mathcal{M}_h^M = \{\mathbf{u}_h(\mu_i) : i = 1, \dots, M\},$$

l'ensemble des solutions associé aux $\{\mu_1, \dots, \mu_M\} \subset \mathcal{P}$.

Interprétation. Pour tout $\mathbf{W} \in \mathcal{V}_n$, on a $P_V = \mathbf{W}\mathbf{W}^T$ pour $V = \text{vectcol}(\mathbf{W})$.

Considérons

$$\mathcal{M}_h^M = \{\mathbf{u}_h(\mu_i) : i = 1, \dots, M\},$$

l'ensemble des solutions associé aux $\{\mu_1, \dots, \mu_M\} \subset \mathcal{P}$. Ainsi (12) implique

$$\sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - P_{V_n} \mathbf{u}_h(\mu_i)\|_2^2 = \min_{V \subset \mathbb{R}^N, \dim V = n} \sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - P_V \mathbf{u}_h(\mu_i)\|_2^2 = \sum_{i=n+1}^r \sigma_i^2.$$

Interprétation. Pour tout $\mathbf{W} \in \mathcal{V}_n$, on a $P_V = \mathbf{W}\mathbf{W}^T$ pour $V = \text{vectcol}(\mathbf{W})$.

Considérons

$$\mathcal{M}_h^M = \{\mathbf{u}_h(\mu_i) : i = 1, \dots, M\},$$

l'ensemble des solutions associé aux $\{\mu_1, \dots, \mu_M\} \subset \mathcal{P}$. Ainsi (12) implique

$$\sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - P_{V_n} \mathbf{u}_h(\mu_i)\|_2^2 = \min_{V \subset \mathbb{R}^N, \dim V = n} \sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - P_V \mathbf{u}_h(\mu_i)\|_2^2 = \sum_{i=n+1}^r \sigma_i^2.$$

Cela signifie que l'espace $V_n = \text{vectcol}(\Phi) = \text{vect}\{\varphi_1, \dots, \varphi_n\}$

1. est **optimal** i.e. le meilleur sous-espace de dimension n de $\mathbb{R}^N$

Interprétation. Pour tout $\mathbf{W} \in \mathcal{V}_n$, on a $P_V = \mathbf{W}\mathbf{W}^T$ pour $V = \text{vectcol}(\mathbf{W})$.

Considérons

$$\mathcal{M}_h^M = \{\mathbf{u}_h(\mu_i) : i = 1, \dots, M\},$$

l'ensemble des solutions associé aux $\{\mu_1, \dots, \mu_M\} \subset \mathcal{P}$. Ainsi (12) implique

$$\sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - P_{V_n} \mathbf{u}_h(\mu_i)\|_2^2 = \min_{V \subset \mathbb{R}^N, \dim V = n} \sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - P_V \mathbf{u}_h(\mu_i)\|_2^2 = \sum_{i=n+1}^r \sigma_i^2.$$

Cela signifie que l'espace $V_n = \text{vectcol}(\Phi) = \text{vect}\{\varphi_1, \dots, \varphi_n\}$

1. est **optimal** i.e. le meilleur sous-espace de dimension n de $\mathbb{R}^N$
2. qui approche $\mathcal{M}_h^M$ en norme 2.

Interprétation. Pour tout $\mathbf{W} \in \mathcal{V}_n$, on a $P_V = \mathbf{W}\mathbf{W}^T$ pour $V = \text{vectcol}(\mathbf{W})$.

Considérons

$$\mathcal{M}_h^M = \{\mathbf{u}_h(\mu_i) : i = 1, \dots, M\},$$

l'ensemble des solutions associé aux $\{\mu_1, \dots, \mu_M\} \subset \mathcal{P}$. Ainsi (12) implique

$$\sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - P_{V_n} \mathbf{u}_h(\mu_i)\|_2^2 = \min_{V \subset \mathbb{R}^N, \dim V = n} \sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - P_V \mathbf{u}_h(\mu_i)\|_2^2 = \sum_{i=n+1}^r \sigma_i^2.$$

Cela signifie que l'espace $V_n = \text{vectcol}(\Phi) = \text{vect}\{\varphi_1, \dots, \varphi_n\}$

1. est **optimal** i.e. le meilleur sous-espace de dimension n de $\mathbb{R}^N$
2. qui approche $\mathcal{M}_h^M$ en norme 2.
3. L'erreur commise quand on considère $V_n \approx \mathcal{M}_h^M$ est $\sum_{i=n+1}^r \sigma_i^2$.

Théorème d'Eckart-Young [QMN, §6.3.1.]

Soit $A \in \mathbb{R}^{N \times M}$. On a

$$\|A - \Phi_n \Phi_n^T A\|_F = \min_{B \in \mathbb{R}^{N \times M}, \text{rang}(B) \leq n} \|A - B\|_F = \sqrt{\sum_{i=n+1}^r \sigma_i^2}.$$

Théorème d'Eckart-Young [QMN, §6.3.1.]

Soit $A \in \mathbb{R}^{N \times M}$. On a

$$\|A - \Phi_n \Phi_n^T A\|_F = \min_{B \in \mathbb{R}^{N \times M}, \text{rang}(B) \leq n} \|A - B\|_F = \sqrt{\sum_{i=n+1}^r \sigma_i^2}.$$

Éléments de preuve (Optimalité). On a

$$\sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - \Phi_n \Phi_n^T \mathbf{u}_h(\mu_i)\|_2^2 = \|A - \Phi_n \Phi_n^T A\|_F^2$$

Théorème d'Eckart-Young [QMN, §6.3.1.]

Soit $A \in \mathbb{R}^{N \times M}$. On a

$$\|A - \Phi_n \Phi_n^T A\|_F = \min_{B \in \mathbb{R}^{N \times M}, \text{rang}(B) \leq n} \|A - B\|_F = \sqrt{\sum_{i=n+1}^r \sigma_i^2}.$$

Eléments de preuve (Optimalité). On a

$$\sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - \Phi_n \Phi_n^T \mathbf{u}_h(\mu_i)\|_2^2 = \|A - \Phi_n \Phi_n^T A\|_F^2$$

Par Eckart Young, on a pour tout B de rang n

$$\|A - \Phi_n \Phi_n^T A\|_F^2 \leq \|A - B\|_F^2$$

Théorème d'Eckart-Young [QMN, §6.3.1.]

Soit $A \in \mathbb{R}^{N \times M}$. On a

$$\|A - \Phi_n \Phi_n^T A\|_F = \min_{B \in \mathbb{R}^{N \times M}, \text{rang}(B) \leq n} \|A - B\|_F = \sqrt{\sum_{i=n+1}^r \sigma_i^2}.$$

Éléments de preuve (Optimalité). On a

$$\sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - \Phi_n \Phi_n^T \mathbf{u}_h(\mu_i)\|_2^2 = \|A - \Phi_n \Phi_n^T A\|_F^2$$

Par Eckart Young, on a pour tout B de rang n

$$\|A - \Phi_n \Phi_n^T A\|_F^2 \leq \|A - B\|_F^2$$

en particulier si $B = \mathbf{W} \mathbf{W}^T A$ avec $W \in \mathcal{V}_n$

$$\|A - \Phi_n \Phi_n^T A\|_F^2 \leq \|A - \mathbf{W} \mathbf{W}^T A\|_F^2$$

Théorème d'Eckart-Young [QMN, §6.3.1.]

Soit $A \in \mathbb{R}^{N \times M}$. On a

$$\|A - \Phi_n \Phi_n^T A\|_F = \min_{B \in \mathbb{R}^{N \times M}, \text{rang}(B) \leq n} \|A - B\|_F = \sqrt{\sum_{i=n+1}^r \sigma_i^2}.$$

Eléments de preuve (Optimalité). On a

$$\sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - \Phi_n \Phi_n^T \mathbf{u}_h(\mu_i)\|_2^2 = \|A - \Phi_n \Phi_n^T A\|_F^2$$

Par Eckart Young, on a pour tout B de rang n

$$\|A - \Phi_n \Phi_n^T A\|_F^2 \leq \|A - B\|_F^2$$

en particulier si $B = \mathbf{W} \mathbf{W}^T A$ avec $W \in \mathcal{V}_n$

$$\|A - \Phi_n \Phi_n^T A\|_F^2 \leq \|A - \mathbf{W} \mathbf{W}^T A\|_F^2$$

On retrouve (12) en passant au min sur $\mathcal{V}_n$.

1. L'efficacité de la SVD est liée aux valeurs singulières de A .
2. Si elles décroissent vite, un espace V_n de petite dimension n suffit pour bien approcher $\mathcal{M}_h^M$!

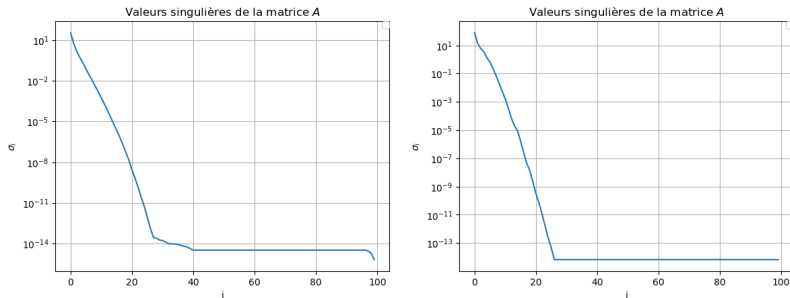


Figure – Valeurs singulières de A pour l'équation d'advection diffusion (gauche) et de la chaleur (droite).

Optimalité en norme 2

$$d_2^M(V_n, \mathcal{M}_h^M) := \left(\sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - P_{V_n} \mathbf{u}_h(\mu_i)\|_2^2 \right)^{1/2}$$

Optimalité en norme 2

$$d_2^M(V_n, \mathcal{M}_h^M) := \left(\sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - P_{V_n} \mathbf{u}_h(\mu_i)\|_2^2 \right)^{1/2}$$

Optimalité en norme ∞

On définit la n -épaisseur de Kolmogorov

$$d_\infty^M(V_n, \mathcal{M}_h^M) := \min_{V \subset \mathbb{R}^N, \dim V = n} \max_{i=1, M} \|\mathbf{u}_h(\mu_i) - P_V \mathbf{u}_h(\mu_i)\|_2 \quad (13)$$

Optimalité en norme 2

$$d_2^M(V_n, \mathcal{M}_h^M) := \left(\sum_{i=1}^M \|\mathbf{u}_h(\mu_i) - P_{V_n} \mathbf{u}_h(\mu_i)\|_2^2 \right)^{1/2}$$

Optimalité en norme ∞

On définit la n -épaisseur de Kolmogorov

$$d_\infty^M(V_n, \mathcal{M}_h^M) := \min_{V \subset \mathbb{R}^N, \dim V = n} \max_{i=1, M} \|\mathbf{u}_h(\mu_i) - P_V \mathbf{u}_h(\mu_i)\|_2 \quad (13)$$

🔍 Par définition des normes 2 et ∞ sur $\mathbb{R}^M$

$$M^{-1/2} d_2^M(V_n, \mathcal{M}_h^M) \leq d_\infty^M(V_n, \mathcal{M}_h^M) \leq d_2^M(V_n, \mathcal{M}_h^M).$$

🔍 Résoudre le problème (13) n'est pas trivial en pratique.

Principe. On construit $\{V_\ell\}_{\ell=0}^n \subset \mathbb{R}^N$ t.q. $V_{\ell-1} \subset V_\ell$ de façon **adaptative**.

Soit $V_0 = \{0\}$. Pour $\ell = 1, \dots, n$ on construit V_ℓ comme il suit.

1) Sélectionner

$$\mu_\ell \in \arg \max_{i=1, \dots, M} \|\mathbf{u}_h(\mu_i) - P_{V_{\ell-1}} \mathbf{u}_h(\mu_i)\|_2.$$

2) Calculer le snapshot $\mathbf{u}_h(\mu_\ell)$ solution de $F_h(\mathbf{u}_h(\mu_\ell), \mu_\ell) = 0$.

3) Poser $V_\ell = \text{vect}\{\mathbf{u}_h(\mu_1), \dots, \mathbf{u}_h(\mu_\ell)\}$.

Principe. On construit $\{V_\ell\}_{\ell=0}^n \subset \mathbb{R}^N$ t.q. $V_{\ell-1} \subset V_\ell$ de façon **adaptative**.

Soit $V_0 = \{0\}$. Pour $\ell = 1, \dots, n$ on construit V_ℓ comme il suit.

1) Sélectionner

$$\mu_\ell \in \arg \max_{i=1, \dots, M} \|\mathbf{u}_h(\mu_i) - P_{V_{\ell-1}} \mathbf{u}_h(\mu_i)\|_2.$$

2) Calculer le snapshot $\mathbf{u}_h(\mu_\ell)$ solution de $F_h(\mathbf{u}_h(\mu_\ell), \mu_\ell) = 0$.

3) Poser $V_\ell = \text{vect}\{\mathbf{u}_h(\mu_1), \dots, \mathbf{u}_h(\mu_\ell)\}$.

✗ Calculer l'erreur de projection

$$\|\mathbf{u}_h(\mu_i) - P_{V_{\ell-1}} \mathbf{u}_h(\mu_i)\|_2$$

nécessite de calculer $\mathbf{u}_h(\mu_i)$ pour tout $i = 1, \dots, M$.

Principe. On construit $\{V_\ell\}_{\ell=0}^n \subset \mathbb{R}^N$ t.q. $V_{\ell-1} \subset V_\ell$ de façon **adaptative**.

Soit $V_0 = \{0\}$. Pour $\ell = 1, \dots, n$ on construit V_ℓ comme il suit.

1) Sélectionner

$$\mu_\ell \in \arg \max_{i=1, \dots, M} \Delta(u_{\ell-1}(\xi), \xi).$$

2) Calculer le snapshot $\mathbf{u}_h(\mu_\ell)$ solution de $F_h(\mathbf{u}_h(\mu_\ell), \mu_\ell) = 0$.

3) Poser $V_\ell = \text{vect}\{\mathbf{u}_h(\mu_1), \dots, \mathbf{u}_h(\mu_\ell)\}$.

✗ Calculer l'erreur de projection

$$\|\mathbf{u}_h(\mu_i) - P_{V_{\ell-1}} \mathbf{u}_h(\mu_i)\|_2$$

nécessite de calculer $\mathbf{u}_h(\mu_i)$ pour tout $i = 1, \dots, M$.

✓ On utilise un estimateur d'erreur e.g.

$$\Delta(u_\ell(\mu_i), \mu_i) := \|L_h(\mu_i) \mathbf{u}_\ell(\mu_i) - \mathbf{f}_h(\mu_i)\|_2$$

ne nécessite pas de connaître $\mathbf{u}_h(\mu_i)$.

Principe. On construit $\{V_\ell\}_{\ell=0}^n \subset \mathbb{R}^N$ t.q. $V_{\ell-1} \subset V_\ell$ de façon **adaptative**.

Soit $V_0 = \{0\}$. Pour $\ell = 1, \dots, n$ on construit V_ℓ comme il suit.

1) Sélectionner

$$\mu_\ell \in \arg \max_{i=1, \dots, M} \Delta(u_{\ell-1}(\xi), \xi).$$

2) Calculer le snapshot $\mathbf{u}_h(\mu_\ell)$ solution de $F_h(\mathbf{u}_h(\mu_\ell), \mu_\ell) = 0$.

3) Poser $V_\ell = \text{vect}\{\mathbf{u}_h(\mu_1), \dots, \mathbf{u}_h(\mu_\ell)\}$.

✗ Calculer l'erreur de projection

$$\|\mathbf{u}_h(\mu_i) - P_{V_{\ell-1}} \mathbf{u}_h(\mu_i)\|_2$$

nécessite de calculer $\mathbf{u}_h(\mu_i)$ pour tout $i = 1, \dots, M$.

✓ On utilise un estimateur d'erreur e.g.

$$\Delta(u_\ell(\mu_i), \mu_i) := \|L_h(\mu_i) \mathbf{u}_\ell(\mu_i) - \mathbf{f}_h(\mu_i)\|_2$$

ne nécessite pas de connaître $\mathbf{u}_h(\mu_i)$.

✓ On a besoin de n snapshots contre M dans la méthode basée sur la SVD.

Les **méthodes d'approximation linéaire** ou **sous format de faible rang**

- ✓ sont efficaces pour des problèmes pour lesquels n -épaisseur de Kolmogorov **décroît rapidement**

ex. : Exponentiellement pour des problèmes de type diffusion

Les méthodes d'approximation linéaire ou sous format de faible rang

- ✓ sont efficaces pour des problèmes pour lesquels n -épaisseur de Kolmogorov décroît rapidement

ex. : Exponentiellement pour des problèmes de type diffusion

- ✗ mais montrent leur limite quand la n -épaisseur de Kolmogorov décroît lentement !

ex. : Comme $n^{-1/2}$ pour les problèmes de type transport

Ca ne marche pas toujours...

Soit $u : [0, T] \times \mathbb{R} \rightarrow \mathbb{R}$ solution de

$$\partial_t u(t, x) + \partial_x u(t, x) = 0, \quad x \in \mathbb{R} \quad (14)$$

avec $u(0, x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{x^2}{2\sigma^2}}$ et $\sigma > 0$. Le problème (14) admet pour solution

$$u(t, x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(x-t)^2}{2\sigma^2}}.$$

Construisons la matrice des snapshots $A \in \mathbb{R}^{N \times M}$ t.q. $A_{ij} = u(t_j, x_i)$.

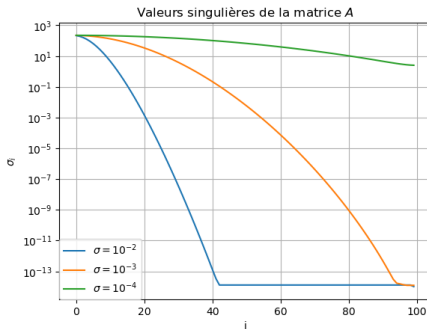
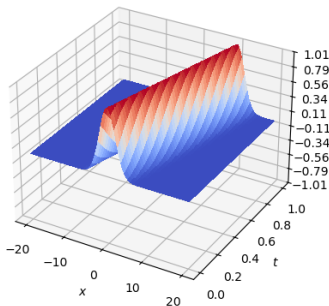


Figure - Gauche : solution u pour $\sigma = 1$. Droite : valeurs singulières de A , $N = 10^4$, $M = 500$.

[QMN] Alfio Quarteroni, Andrea Manzoni et Federico Negri. Reduced Basis Methods for Partial Differential Equations. Cham, Switzerland : Springer International Publishing (cf. p. 91, 92, 99-103).