

High Performance Computing Uncertainty Quantification (Overview of Mathematical Tools)

P. Dossantos-Uzarralde

Option SGI



Uncertainty Sources

Epistemic Uncertainty

- Lack of knowledge about, say, the appropriate value to use for a quantity, or the proper model form to use.
"Reducible uncertainty" : can be reduced through increased understanding (research) or more, relevant data.

Aleatory Uncertainty

- "Alea" = Latin for "die" ; Latin aleator = "dice player."
Inherent randomness, intrinsic variability.
"Irreducible uncertainty" : cannot be reduced by additional data.
Usually modeled with probability distributions.

Uncertainty Quantification : Some facts

- UQ in computational science is the formal characterization, propagation, aggregation, comprehension, and communication of aleatory (variability) and epistemic (incomplete knowledge) uncertainties.
 - E.g., demand fluctuations in a power grid are aleatory uncertainties.
 - E.g., incomplete knowledge about the future (scenarios uncertainty), the validity of models, and inadequate "statistics" are epistemic uncertainties.

- A huge range of technical issues arise in the Modeling & Simulation (M&S) components of problem definition and execution phases.

- Another huge range of technical issues arises in the delivery phase, especially in high-risk decision environments.

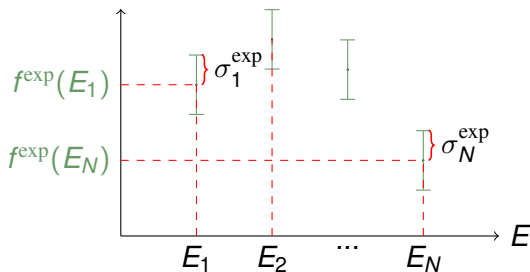
- "Probability" is the main foundation for current "quantification".

UQ impact on HPC

- Any large-scale computational problem's computing requirements increase (usually significantly) with UQ.

Problem

N measurements $f^{\text{exp}} = (f^{\text{exp}}(E_1), \dots, f^{\text{exp}}(E_N))$

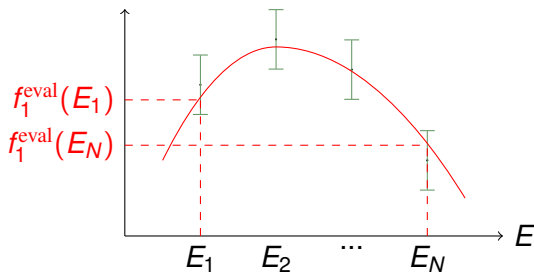


Problem

One possible model

N measurements $f^{\text{exp}} = (f^{\text{exp}}(E_1), \dots, f^{\text{exp}}(E_N))$

First set of input parameters: $p^1 = (p^1_1(E_1), \dots, p^1_{n_p}(E_N))$

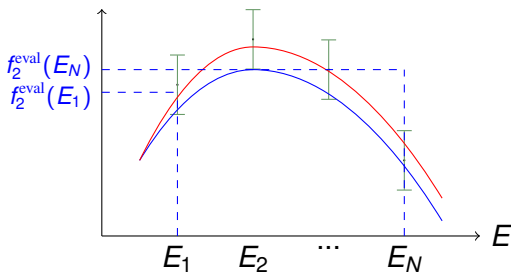


Problem

Two possible models

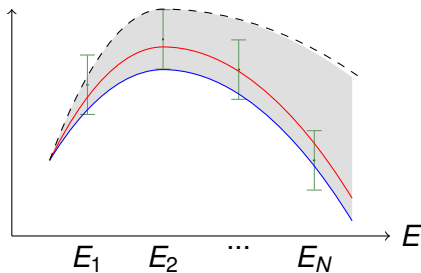
N measurements $f^{\text{exp}} = (f^{\text{exp}}(E_1), \dots, f^{\text{exp}}(E_N))$

Second set of input parameters: $p^2 = (p_1^2(E_1), \dots, p_{n_p}^2(E_N))$



Problem

N measurements $f^{\text{exp}} = (f^{\text{exp}}(E_1), \dots, f^{\text{exp}}(E_N))$
 n sets of input parameters $\Rightarrow$ n possible models

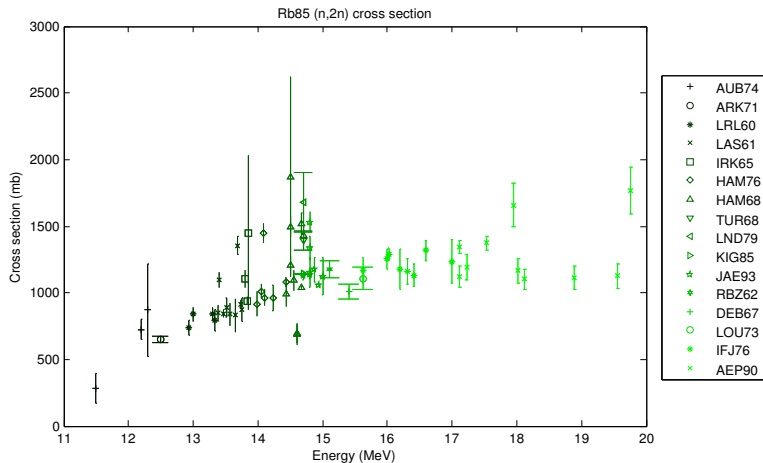


"Degrees of Goodness"

- When we use experimental results (such as property values) in an analytical solution, we should consider "how good" the data are and what influence that degree of goodness has on the interpretation and usefulness of the solution
- When we compare model predictions with experimental data, as in a validation process, we should consider the degree of goodness of the model results and the degree of goodness of the data.

Transposition to the true life ?

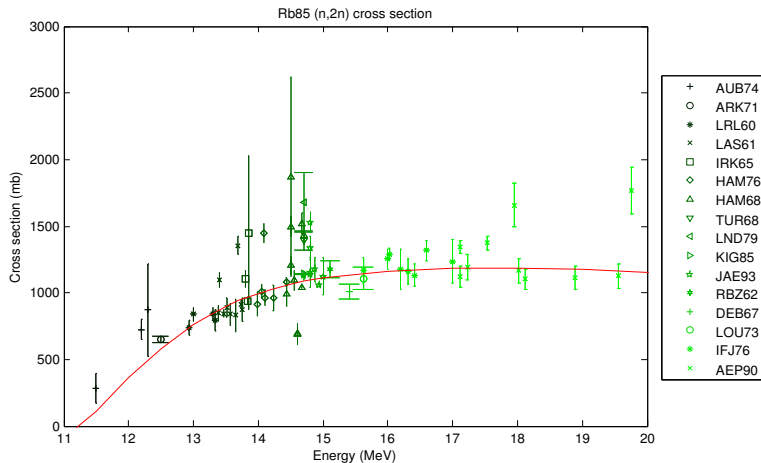
Real Physical Problem



Evaluation of nuclear cross sections.

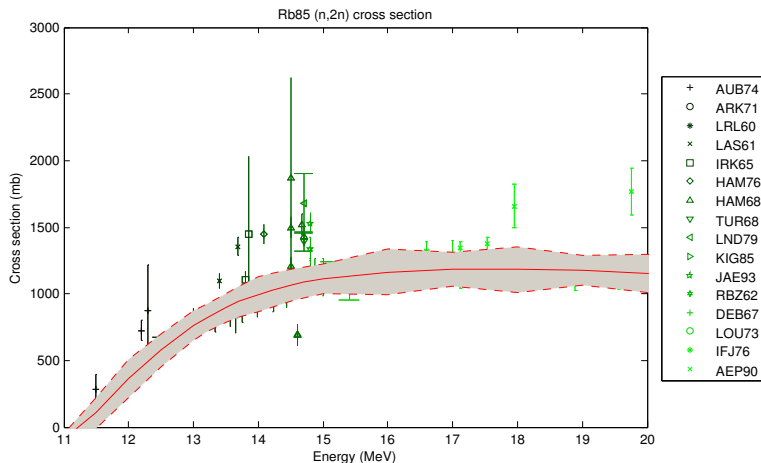
Experimental Data .

Context: evaluation



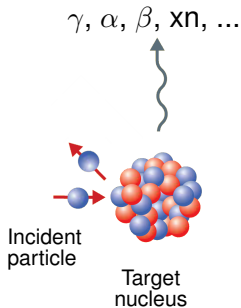
Find the good analytical model

Context: evaluation and covariance matrix



Find the model with the associated uncertainties estimations

What is the model ?



Shrödinger's Equation

Cross sections Evaluation

$$\nabla^2 \psi(\vec{r}) + \frac{2 \cdot m}{\hbar^2} \cdot (E - \mathcal{U}(r, E)) \cdot \psi(\vec{r}) = 0$$

Diagram illustrating the components of the Schrödinger equation:

- ψ : wave function of the incident particle
- $\vec{r}$: position
- m : the reduced mass of the system {neutron + nucleus}
- E : energy
- $\mathcal{U}(r, E)$: potential created by the nucleus
- r : radial coordinate of $\vec{r}$

Potentiel Optique

$$\mathcal{U}(r, E) = \underbrace{-\mathcal{V}_D(r, E) - i\mathcal{W}_D(r, E)}_{\text{volume}} \underbrace{-\mathcal{V}_S(r, E) - i\mathcal{W}_S(r, E)}_{\text{surface}} \underbrace{+\mathcal{V}_{SO}(r, E) + i\mathcal{W}_{SO}(r, E)}_{\text{spin-orbite}}$$

volume

surface

spin-orbite

$$\forall j \in \{D, S, SO\}$$

$$\mathcal{V}_j(r, E) = V_j(E) \cdot g(r, R_{Vj}, a_{Vj}) \quad \text{and} \quad \mathcal{W}_j(r, E) = W_j(E) \cdot g(r, R_{Wj}, a_{Wj})$$

Paramètres du modèle optique

Potentiel Optique

$$\mathcal{U}(r, E) = \underbrace{-\mathcal{V}_D(r, E) - i\mathcal{W}_D(r, E)}_{\text{volume}} \underbrace{-\mathcal{V}_S(r, E) - i\mathcal{W}_S(r, E)}_{\text{surface}} \underbrace{+\mathcal{V}_{SO}(r, E) + i\mathcal{W}_{SO}(r, E)}_{\text{spin-orbite}}$$

volume

surface

spin-orbite

$$\forall j \in \{D, S, SO\}$$

$$\mathcal{V}_j(r, E) = V_j(E) \cdot g(r, R_{Vj}, a_{Vj}) \quad \text{and} \quad \mathcal{W}_j(r, E) = W_j(E) \cdot g(r, R_{Wj}, a_{Wj})$$

Paramètres du modèle optique

18 paramètres: $p(E)$

Objectives

Introduce

- Challenges in UQ Science
- Various mathematical techniques

Convey

- To advance the application of mathematics and computational science to engineering, industry, science, and society
- To promote research that will lead to effective new mathematical and computational methods and techniques for science, engineering, industry, and society

Outline

- 1 Taylor Series Method
- 2 Monte Carlo Method
- 3 Global Sensitivity Analysis
 - Sobol indices
 - Introduction to Chaos Polynomial
 - Quasi-Monte Carlo Application
- 4 Least-Square vs Bayesian inference
- 5 Data learning application
- 6 Metric on matrix space
- 7 Bootstrap utility

Outline

- 1 Taylor Series Method
- 2 Monte Carlo Method
- 3 Global Sensitivity Analysis
 - Sobol indices
 - Introduction to Chaos Polynomial
 - Quasi-Monte Carlo Application
- 4 Least-Square vs Bayesian inference
- 5 Data learning application
- 6 Metric on matrix space
- 7 Bootstrap utility

First order Taylor's development.

The Taylor series method is to approximate f_j by a linear function. The linearization greatly simplifies the error analysis, but at the expense of introducing an approximation error.

$$Y_j \approx f_j(\mu_1^X, \dots, \mu_p^X) + \sum_{i=1}^p \left(\frac{\partial f_j}{\partial X_i}(\mu_1^X, \dots, \mu_p^X) \right) (X_i - \mu_i^X). \quad (1)$$

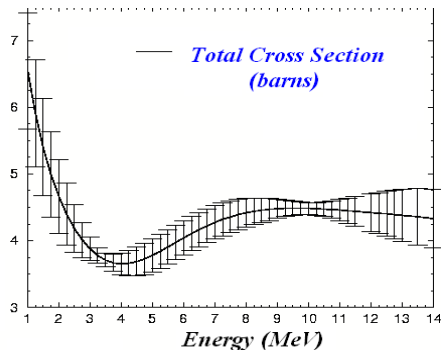
- Mean Value, Variance, Covariance :

$$\mu_j^Y \approx f_j(\mu_1^X, \dots, \mu_p^X), \quad \sigma_{jj}^Y \approx \sum_i \sum_{i'} \frac{\partial f_j}{\partial X_i} \frac{\partial f_j}{\partial X_{i'}} \sigma_{ii'}^X, \quad \sigma_{jj'}^Y \approx \sum_i \sum_{i'} \frac{\partial f_j}{\partial X_i} \frac{\partial f_{j'}}{\partial X_{i'}} \sigma_{ii'}^X.$$

In a matrix form (Sandwich rule Cacuci[1981]) as : $\underline{\underline{V_Y}} = \underline{\underline{F_X}} \underline{\underline{V_X}} \underline{\underline{F_X}}^T$,
 where $\underline{\underline{V_X}}$ here denotes the diagonal covariance matrix and $\underline{\underline{F_X}}$ the sensitivity matrix.

First order Taylor's development.

- Results



Second order Taylor's development

Goal : estimate how the second order affects the sensitivity of the results

$$\begin{aligned}
 Y_j &\approx f_j(\mu_1^X, \dots, \mu_p^X) + \sum_{i=1}^p \left(\frac{\partial f_j}{\partial X_i}(\mu_1^X, \dots, \mu_p^X) \right) (X_i - \mu_i^X) \\
 &+ \frac{1}{2} \sum_{i=1}^p \sum_{i'=1}^p \left(\frac{\partial^2 f_j}{\partial X_i \partial X_{i'}}(\mu_1^X, \dots, \mu_p^X) \right) (X_i - \mu_i^X)(X_{i'} - \mu_{i'}^X)
 \end{aligned} \quad (2)$$

Under the assumptions of independent parameters and continuous uniform distribution, the final covariance $\text{Cov}(Y_j, Y_{j'}) = \sigma_{jj'}^Y$ can be written:

$$\sigma_{jj'}^Y = \sum_{i=1}^p \frac{\partial f_j}{\partial X_i} \frac{\partial f_{j'}}{\partial X_i} \sigma_{ii}^X + \frac{1}{4} \sum_{i=1}^p \frac{\partial^2 f_j}{\partial X_i^2} \frac{\partial^2 f_{j'}}{\partial X_i^2} \left(E[(X_i - \mu_i^X)^4] - (\sigma_{ii}^X)^2 \right).$$

Introducing the second derivatives and the 4th central moment diagonal matrix:

$$\underline{\underline{F_X^2}} = \left(\frac{\partial^2 f_j}{\partial X_i^2} \right)_{\substack{1 \leq j \leq g \\ 1 \leq i \leq p}} \quad \underline{\underline{Q_X}} = (Q_{Xij})_{\substack{1 \leq j \leq p \\ 1 \leq i \leq p}} = \begin{cases} 0 & \text{si } i \neq j \\ E[(X_i - \mu_i^X)^4] & \text{si } i = j \end{cases}$$

The results above can be written in a matrix form as:

$$\underline{\underline{F_Y}} = \underline{\underline{F_X}} \underline{\underline{V_X}} \underline{\underline{F_X}}^T + \frac{1}{4} \underline{\underline{F_X^2}} (\underline{\underline{Q_X}} - \underline{\underline{V_X}}^2) \underline{\underline{F_X}}^T \quad (3)$$

Outline

- 1 Taylor Series Method
- 2 Monte Carlo Method**
- 3 Global Sensitivity Analysis
 - Sobol indices
 - Introduction to Chaos Polynomial
 - Quasi-Monte Carlo Application
- 4 Least-Square vs Bayesian inference
- 5 Data learning application
- 6 Metric on matrix space
- 7 Bootstrap utility

Monte Carlo Method

Monte Carlo simulation is categorized as a sampling method because the inputs are randomly generated from probability distributions.

- Uniform distribution

In the context of a comparison between both methods, we used the uniform distribution $\mathcal{U}(\alpha_i, \beta_i)$ as parameter probability function distribution (pdf).

$$f_{X_i}(x) = \frac{1}{\beta_i - \alpha_i} \mathbb{I}_{[\alpha_i, \beta_i]}(x), \quad \mathbb{I}_{[\alpha_i, \beta_i]}(x) = \begin{cases} 0 & \text{si } x \notin [\alpha_i, \beta_i] \\ 1 & \text{si } x \in [\alpha_i, \beta_i] \end{cases}$$

- Independent Parameters

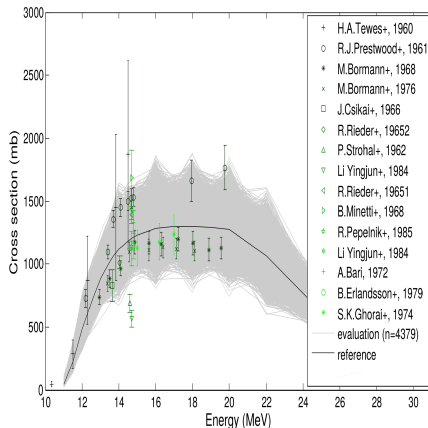
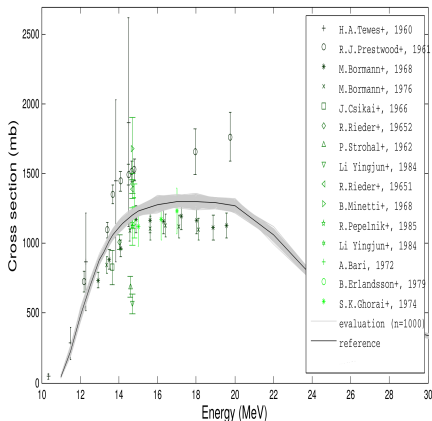
The pdf of the parameter vector $\underline{X}$ is the product of the marginal distributions f_{X_i} of each parameters X_i :

$$f_{\underline{X}} = \prod_{i=1}^p f_{X_i}.$$

Simulations Monte-Carlo

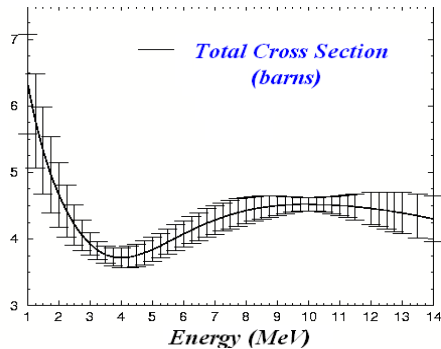
Parameters :

$r_V, a_V, v_1, v_2, v_3, w_1, w_2, r_D, a_D, d_1, d_2, d_3, r_{SO}, a_{SO}, v_{SO_1}, v_{SO_2}, w_{SO_1}$ et w_{SO_2}

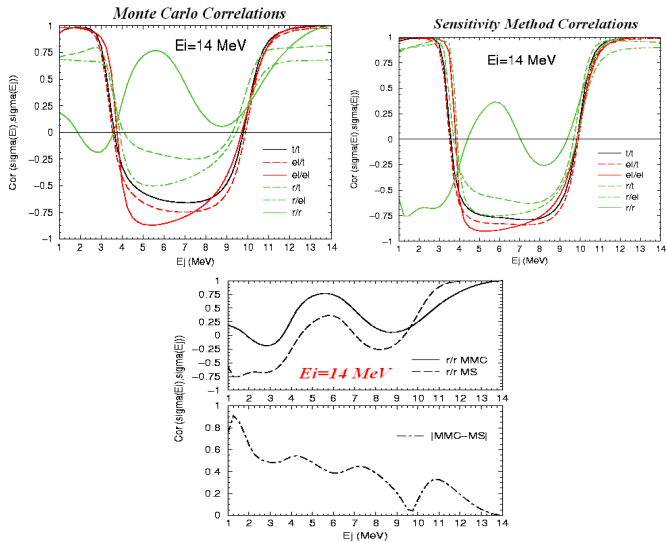


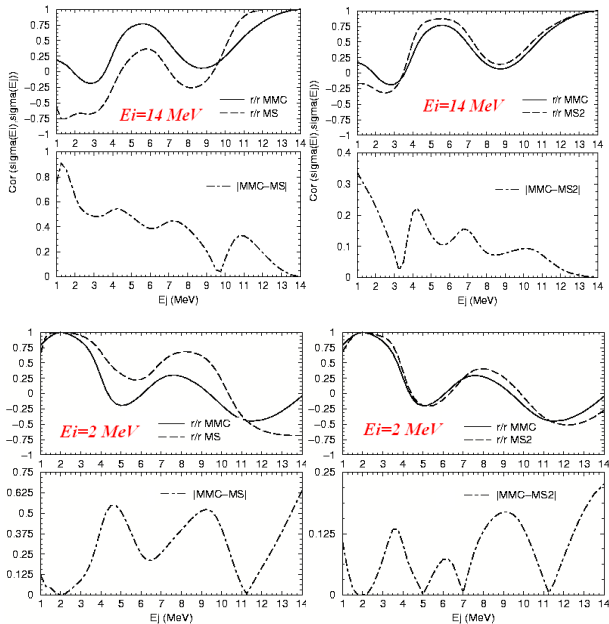
Monte Carlo Method

- Results



Comparisons





Local Sensitivity Analysis

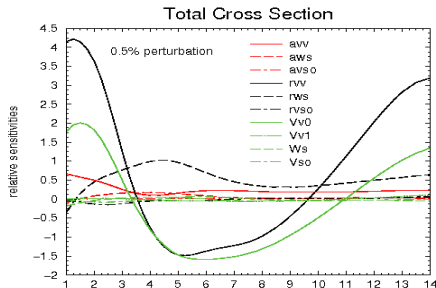
The aim of sensitivity analysis : estimate the rate of change in the model output with respect to changes in model inputs - determine parameters for which it is important to have more accurate values, and understanding the behavior of the system being modeled.

Pb : the method involves an expansion of model outputs in terms of small random perturbations of model parameters.

- Different δx_i were tested in the interval $[0, \alpha x_i = \pm 5\%]$.
- In this study, we use the normalized gradient $\frac{\partial f_j}{\partial x_i} \frac{\mu_j^x}{f_j}$.

Main limitation : require the perturbation terms be small.

These methods are in general difficult to apply along with the modeling of complex, nonlinear systems.



Outline

- 1 Taylor Series Method
- 2 Monte Carlo Method
- 3 Global Sensitivity Analysis**
 - Sobol indices
 - Introduction to Chaos Polynomial
 - Quasi-Monte Carlo Application
- 4 Least-Square vs Bayesian inference
- 5 Data learning application
- 6 Metric on matrix space
- 7 Bootstrap utility

Sensitivity Analysis

Purpose of Sensitivity Analysis (SA)

Sensitivity analysis is the study of how the uncertainty in the output of a mathematical model or system (numerical or otherwise) can be apportioned to different sources of uncertainty in its inputs. The process of recalculating outcomes under alternative assumptions to determine the impact of variable under sensitivity analysis can be useful for a range of purposes including :

- Testing the robustness of the results of a model or system in the presence of uncertainty.
- Increased understanding of the relationships between input and output variables in a system or model.
- Identifying model inputs that cause significant uncertainty in the output.
- Searching for errors in the model (by encountering unexpected relationships between inputs and outputs).
- Model simplification - fixing model inputs that have no effect on the output, or identifying and removing redundant parts of the model structure.
- Finding regions in the space of input factors for which the model output is either maximum or minimum or meets some optimum criterion (see optimization and Monte Carlo filtering).

Sobol Indices

- Let us assume the model function $y = f(x_1, \dots, x_p)$, where $x_1, \dots, x_p$ are independent input factors. Whenever $y = f(x_1, x_2, \dots, x_p)$ is integrable over $[0, 1]^p$, y can be decomposed (Sobol[1993]) as :

$$f(x_1, \dots, x_p) = f_0 + \sum_{i=1}^p f_i(x_i) + \sum_{1 \leq i < j \leq p} f_{ij}(x_i, x_j) + \dots + f_{1, \dots, p}(x_1, \dots, x_p)$$

with the classical property of orthogonality.

- Consider that the input parameters are independent random variables, the joint pdf of the input factors is $P(X_1, X_2, \dots, X_p) = \prod_{i=1}^p p(X_i)$.

$$E[f(X_1, \dots, X_p)] = \int_{[0,1]^p} \left(f_0 + \sum_{i=1}^p f_i(x_i) + \sum_{1 \leq i < j \leq p} f_{ij}(x_i, x_j) + \dots + f_{12\dots p}(x_1, x_2, \dots, x_p) \right) \prod_{i=1}^p p(x_i) dx_i$$

$$V(Y) = V(f(X_1, \dots, X_k)) = \int_{[0,1]^p} f^2(X_1, \dots, X_k) \prod_{i=1}^p p(x_i) dx_i - f_0^2$$

- First order Sobol index (S_i) gives the influence of each parameter taken alone.
Second order Sobol index (S_{ij}) a sensitivity measure due to the uncertainties in the set of input parameters $\{i, j\}$.
Higher orders give mixed influence of parameters

$$S_i = \frac{V(E(Y|X_i))}{V(Y)}, \quad S_{ij} = \frac{V(E(Y|X_i, X_j)) - V(E(Y|X_i)) - V(E(Y|X_j))}{V(Y)}, \quad S_{ijk}, \dots$$

Wiener-Askey Polynomial Chaos

The chaos expansion expresses the random process through a complete and orthogonal basis in terms of random variables. Polynomial Chaos {PC} is restricted to second-order stochastic processes, i.e. processes with finite second-order moments (Cameron, Martin[1947]).

- A second-order random process ($\mathcal{L}^2$) can be represented as :

$$\theta(\omega) = \sum_{k=0}^{\infty} \theta_k \Psi_k(\xi(\omega)), \quad \xi(\omega) = \{\xi_i(\omega)\}_{i=1}^{\infty} \quad \text{Wiener[1938]}$$

- θ_k deterministic expansion coefficients,
- $\Psi_k(\xi(\omega))$ random trial basis,
- Ψ'_k 's orthogonal polynomials from the Askey-scheme,
- $\{\xi_i(\omega)\}_{i=1}^{\infty}$ independent random variables.
- Generalized {PC} forms a complete orthonormal basis of the Hilbert space with the inner product and the orthogonality relation.
- In practical computations, the stochastic space must be truncated.
The total number of expansion terms $P + 1$ is determined by N ($\xi(\omega) = \{\xi_i(\omega)\}_{i=1}^N$) and p , which is the highest degree of the orthogonal polynomials :

$$\theta(\omega) = \sum_{k=0}^P \theta_k \Psi_k(\xi(\omega)) \quad P + 1 = \frac{(N + p)!}{N!p!}$$

- The combination of random vector and polynomials is carefully selected from the distribution of the random input.

Sobol Indices Calculations

It is possible to approximately represent the random response of a system as a {PC} expansion. Due to the orthogonality of the basis we can obtain the mean and the variance of the response

$$E[f(X)] = E\left[\sum_{k=0}^{\infty} f_k \psi_k(X)\right] = f_0, \quad V[f(X)] = E\left[\left(\sum_{k=0}^{\infty} f_k \psi_k(X) - f_0\right)^2\right] = \sum_{k=1}^{\infty} f_k^2 E[\psi_k^2(X)]$$

- Expansion coefficients $\{f_k\}_{k=1}^{P-1}$ calculations:

$$\langle f(X) \psi_k(X) \rangle = \left\langle \sum_{l=0}^{\infty} f_l \psi_l(X) \psi_k(X) \right\rangle = \sum_{l=0}^{\infty} f_l \langle \psi_l(X) \psi_k(X) \rangle = f_k \langle \psi_k^2(X) \rangle$$

$$f_k = \frac{\langle f(X) \psi_k(X) \rangle}{\langle \psi_k^2(X) \rangle}$$

- Variance calculation

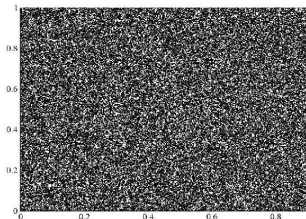
$$\text{Var}[f(X_{i_1}, \dots, X_{i_s})] = E\left[\left(\sum_{k \in \Gamma_{i_1, \dots, i_s}} f_k \psi_k(X)\right)^2\right] = \sum_{k \in \Gamma_{i_1, \dots, i_s}} f_k^2 E[\psi_k^2(X)]$$

- Sobol Indices Calculations** : from the {PC} chaos representation we can reach the full list of Sobol indices (Sudret[2006])

$$S_{i_1, \dots, i_s} = \frac{\sum_{k \in \Gamma_{i_1, \dots, i_s}} f_k^2 E[\psi_k^2(X)]}{\sum_{k=1}^{\infty} f_k^2 E[\psi_k^2(X)]}$$

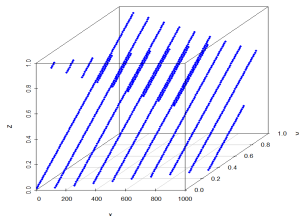
Quasi-Monte Carlo

Congruential Generator

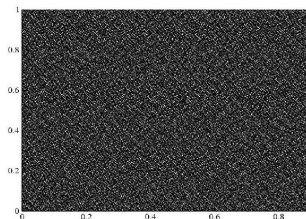


Curse of dimensionality

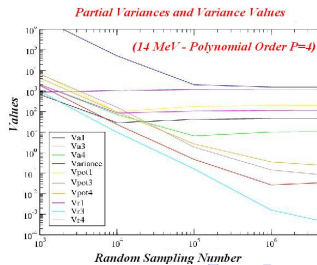
Suite de Halton multidimensionnée



Halton Sequence



Rate of convergence



Outline

- 1 Taylor Series Method
- 2 Monte Carlo Method
- 3 Global Sensitivity Analysis
 - Sobol indices
 - Introduction to Chaos Polynomial
 - Quasi-Monte Carlo Application
- 4 Least-Square vs Bayesian inference**
- 5 Data learning application
- 6 Metric on matrix space
- 7 Bootstrap utility

Least Square

- The most probable probability density function is the Gaussian multivariate distribution :

$$f_{\underline{X}}(\underline{X}) = \frac{1}{(2\pi)^{p/2} |\underline{V}_X|^{1/2}} \exp \left(-\frac{1}{2} (\underline{X} - \underline{\mu})^T \underline{V}_X^{-1} (\underline{X} - \underline{\mu}) \right)$$

- $\underline{V}_X$: in a rigorous way, we introduce the parameters correlations

- $\underline{V}_X$ obtained with a minimum chi-square estimation.

$$\chi^2 = \frac{1}{N} \sum_{i=1}^N \left(\frac{\sigma_i^{exp} - \sigma_i^{calc}}{\Delta \sigma_i^{exp}} \right)^2.$$

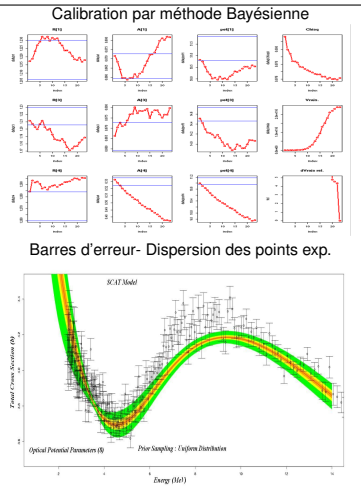
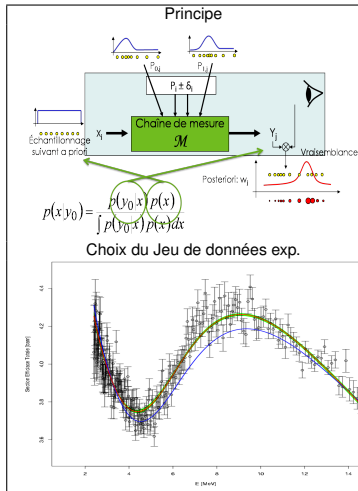
- $\underline{V}$: experimental covariances matrix

- σ_i^{exp} : no covariance estimation.
- introduce χ^2 (generalized χ^2) with experimental covariances :

$$\chi^2 = \sum_{i,j=1}^N \left((\sigma_i^{exp} - \sigma_i^{calc}) (\underline{V}^{-1})_{ij} (\sigma_j^{exp} - \sigma_j^{calc}) \right)$$

- Levenberg-Marquardt Minimization Algorithm (P. Dossantos-Uzarralde[2007])

Bayes inference



● On the right, Bayesian Simulations.

Green : standard error obtained by least square optimisations - Brown : Bayesian optimisations - random parameter interval $\pm 1\%$, in yellow: random parameter interval $\pm 0.5\%$.)

Outline

- 1 Taylor Series Method
- 2 Monte Carlo Method
- 3 Global Sensitivity Analysis
 - Sobol indices
 - Introduction to Chaos Polynomial
 - Quasi-Monte Carlo Application
- 4 Least-Square vs Bayesian inference
- 5 Data learning application**
- 6 Metric on matrix space
- 7 Bootstrap utility

Support Vector Machine

$\omega = (\omega_1, \dots, \omega_d) \in \mathbb{R}^d$ un vecteur de paramètres du modèle f^{model}
 ω_0 le paramètre d'ordonnée à l'origine.

Problème de régression

On cherche une application linéaire $f^{model}(E^{pred}) = {}^t\omega E^{pred} + \omega_0$ telle que pour tout point d'apprentissage de $\mathcal{A}_f$, σ_i^{exp} soit le plus proche possible de $f^{model}(E_i^{exp})$.

E_i^{exp} ($E_i^{exp} \in \mathbb{R}$) : vecteur caractérisant les valeurs des énergies. On pose cependant $E_i^{exp} \in \mathbb{R}^d$ uniquement pour des besoins mathématiques.
 $\langle ., . \rangle$ est le produit scalaire dans $\mathbb{R}^d$.

Support Vector Machine

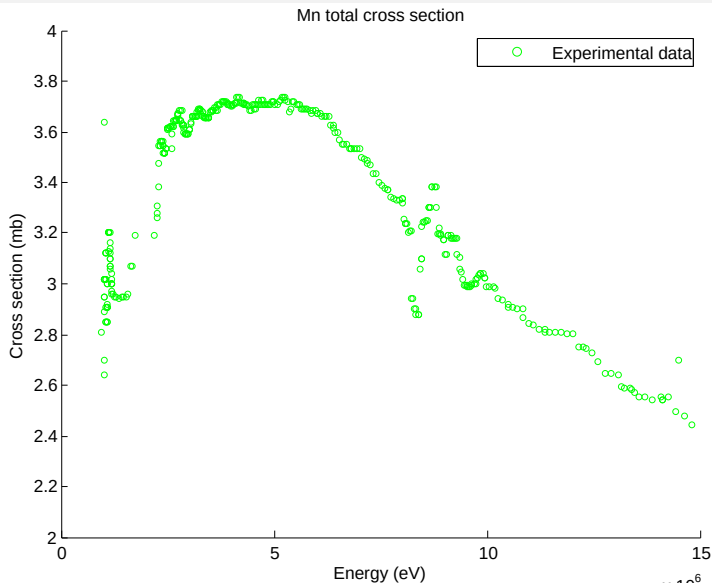
Fonction de perte

$$\begin{aligned} \forall \epsilon \in \mathbb{R}_+, l(\sigma_i^{exp}, f^{model}(E_i^{exp})) &= |\sigma_i^{exp} - f^{model}(E_i^{exp})|_{\epsilon} \\ &= \begin{cases} 0 & \text{si } |\sigma_i^{exp} - f^{model}(E_i^{exp})| \leq \epsilon \\ |\sigma_i^{exp} - f^{model}(E_i^{exp})| - \epsilon & \text{sinon} \end{cases} \end{aligned}$$

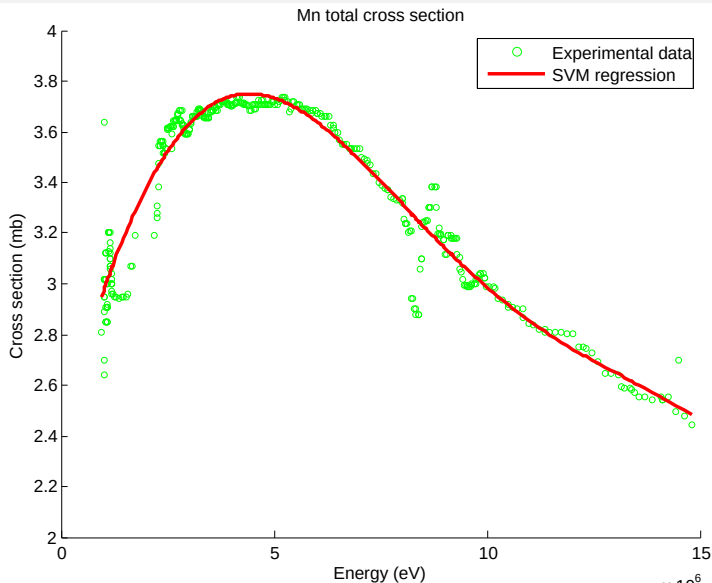
Cette fonction de perte est une norme L_1 epsilon insensible, c'est-à-dire qu'elle ne pénalise pas les mauvaises prédictions à epsilon près. Toutes les valeurs présentes dans le tube d'epsilon sont alors déterminantes pour la régression, les autres valeurs sont pénalisées. En utilisant cette fonction de perte, on obtient le risque réel empirique suivant :

$$\hat{\mathcal{R}}_{Reel}(f^{model}, \mathcal{A}_f) = \frac{1}{m} \sum_{i=1}^m l(\sigma_i^{exp}, f_i^{model}) = \frac{1}{m} \sum_{i=1}^m |\sigma_i^{exp} - f^{model}(E_i^{pred})|_{\epsilon}. \quad (4)$$

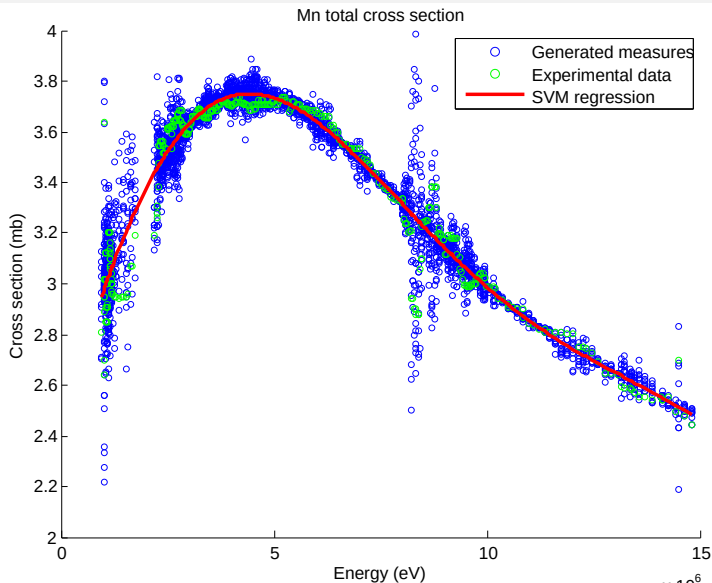
Support Vector Machine



Support Vector Machine

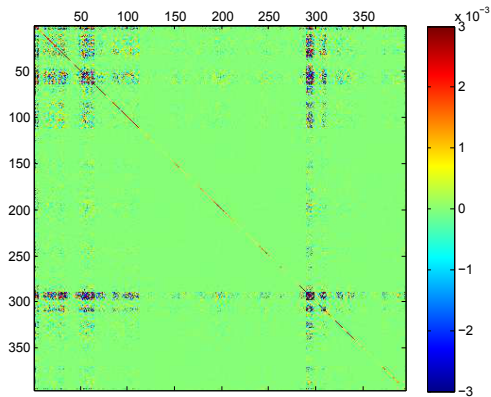


Support Vector Machine

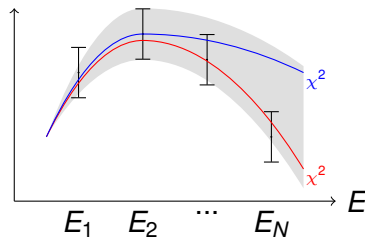
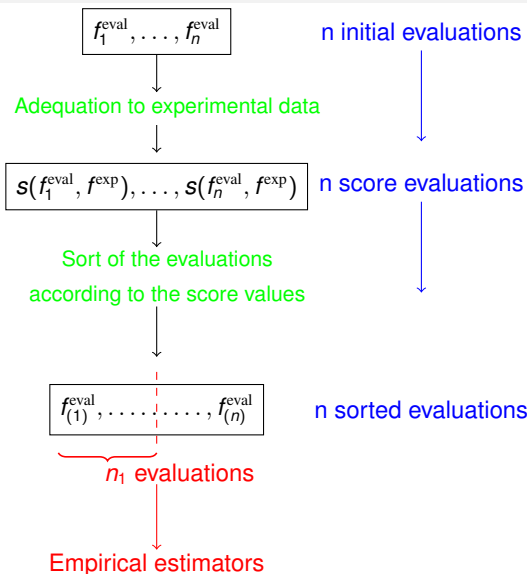


Support Vector Machine

Matrice de covariance "expérimentale" Σ_F
(400, 400)



Scoring Method



$$\chi^2 = \sum_{j=1}^N \left(\frac{f^{\text{eval}}(E_j) - f^{\text{exp}}(E_j)}{\sigma_j^{\text{exp}}} \right)^2$$

Score functions used

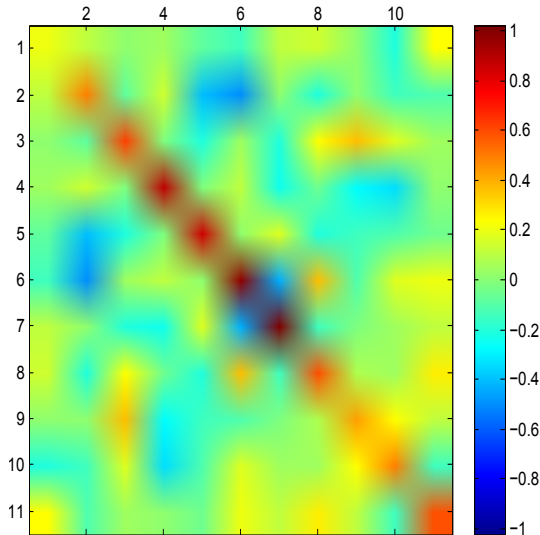
$$\chi^2 \quad S(f_i^{\text{eval}}, f^{\text{exp}}) = \sum_{j=1}^N \left(\frac{f_i^{\text{eval}}(E_j) - f^{\text{exp}}(E_j)}{\sigma_j^{\text{exp}}} \right)^2$$

$$\text{BY} \quad S(f_i^{\text{eval}}, f^{\text{exp}}) = \sum_{j=1}^N \left(\frac{f_i^{\text{eval}}(E_j) - f^{\text{exp}}(E_j)}{\sigma_j^{\text{exp}}} \right)^2 + \sum_{k=1}^{n_p} \sum_{j=1}^N \left(\frac{p_k^j(E_j) - \overline{p_k(E_j)}}{\sigma_k^{\text{param}}(E_j)} \right)^2$$

$$\text{HRBC} \quad S(f_i^{\text{eval}}, f^{\text{exp}}) = \sum_{j=1}^N \frac{\mathbf{I}_{[f^{\text{exp}}(E_j) - \sigma_j^{\text{exp}}; f^{\text{exp}}(E_j) + \sigma_j^{\text{exp}}]}(f_i^{\text{eval}}(E_j))}{\sigma_j^{\text{exp}} \sqrt{2\pi}} \exp \left[-\frac{1}{2} \left(\frac{f_i^{\text{eval}}(E_j) - f^{\text{exp}}(E_j)}{\sigma_j^{\text{exp}}} \right)^2 \right]$$

$$\begin{aligned} \text{with} \quad \mathbf{I}_A(x) &= 1 \text{ if } x \in A \\ \mathbf{I}_A(x) &= 0 \text{ if } x \notin A \end{aligned}$$

Example : Covariance Matrice Estimator



Outline

- 1 Taylor Series Method
- 2 Monte Carlo Method
- 3 Global Sensitivity Analysis
 - Sobol indices
 - Introduction to Chaos Polynomial
 - Quasi-Monte Carlo Application
- 4 Least-Square vs Bayesian inference
- 5 Data learning application
- 6 Metric on matrix space**
- 7 Bootstrap utility

Distance between matrices

Classical distance:
Frobenius norm

Quality of an estimation $\hat{\Sigma}$ of Σ

$$\blacksquare d_F(\hat{\Sigma}, \Sigma) = \|\hat{\Sigma} - \Sigma\|_F = \sqrt{\text{trace}((\hat{\Sigma} - \Sigma)^t \cdot (\hat{\Sigma} - \Sigma))}$$

Quality of the inverse $\hat{\Sigma}^{-1}$ for the estimation of Σ^{-1}

$$\blacksquare d_F(\hat{\Sigma}^{-1}, \Sigma^{-1}) = \|\hat{\Sigma}^{-1} - \Sigma^{-1}\|_F$$

Distance between matrices

Entropy and Kullback-Leibler

Quality of an estimation $\hat{\Sigma}$ of Σ

- Entropy: $d_E(\hat{\Sigma}, \Sigma) = \text{tr}(\hat{\Sigma}.\Sigma^{-1}) - \log(\det(\hat{\Sigma}.\Sigma^{-1})) - N$

Quality of the inverse $\hat{\Sigma}^{-1}$ for the estimation of Σ^{-1}

- Kullback-Leibler: $d_K(\hat{\Sigma}, \Sigma) = \text{tr}(\hat{\Sigma}^{-1}.\Sigma) - \log(\det(\hat{\Sigma}^{-1}.\Sigma)) - N$

James and Stein 1961, *Estimation with quadratic loss*, Proceedings of the Fourth Berkeley Symposium on Mathematical Statistics and Probability.

Levina & al. 2008, *Sparse estimation of large covariance matrices via a nested lasso penalty*, The Annals of Applied Statistics.

Tumminello & al. 2007, *Kullback-Leibler distance as a measure of the information filtered from multivariate data*, Phys. Review E.

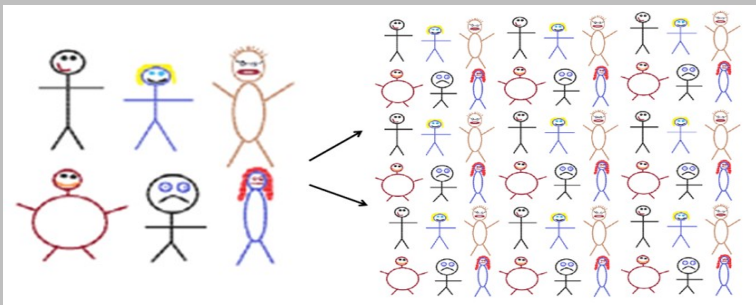
Outline

- 1 Taylor Series Method
- 2 Monte Carlo Method
- 3 Global Sensitivity Analysis
 - Sobol indices
 - Introduction to Chaos Polynomial
 - Quasi-Monte Carlo Application
- 4 Least-Square vs Bayesian inference
- 5 Data learning application
- 6 Metric on matrix space
- 7 Bootstrap utility**

Bootstrap definition

What is Bootstrap

In statistics, bootstrapping can refer to any test or metric that relies on random sampling with replacement. Bootstrapping allows assigning measures of accuracy (defined in terms of bias, variance, confidence intervals, prediction error or some other such measure) to sample estimates.



Bootstrap: general principle

$$\left. \begin{array}{l} X \sim F \\ \Sigma = T(F) \end{array} \right\} \begin{array}{l} \text{what we want} \\ \text{unknown} \end{array}$$

Bootstrap: general principle

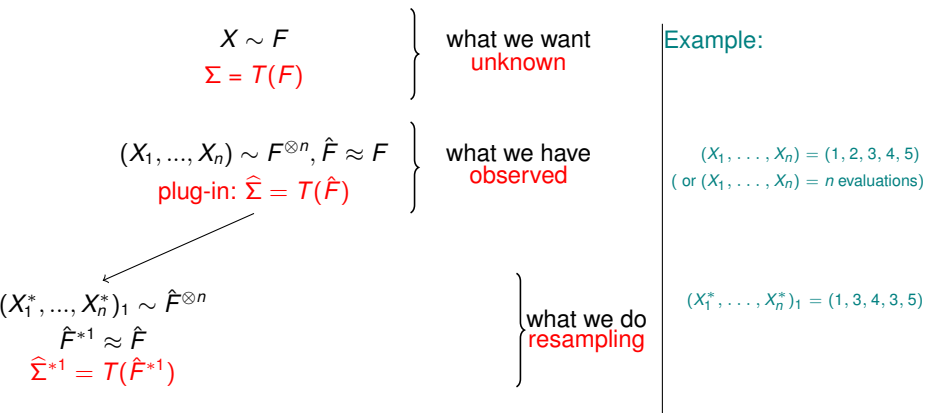
$$\left. \begin{array}{l} X \sim F \\ \Sigma = T(F) \end{array} \right\} \begin{array}{l} \text{what we want} \\ \text{unknown} \end{array}$$

$$\left. \begin{array}{l} (X_1, \dots, X_n) \sim F^{\otimes n}, \hat{F} \approx F \\ \text{plug-in: } \hat{\Sigma} = T(\hat{F}) \end{array} \right\} \begin{array}{l} \text{what we have} \\ \text{observed} \end{array}$$

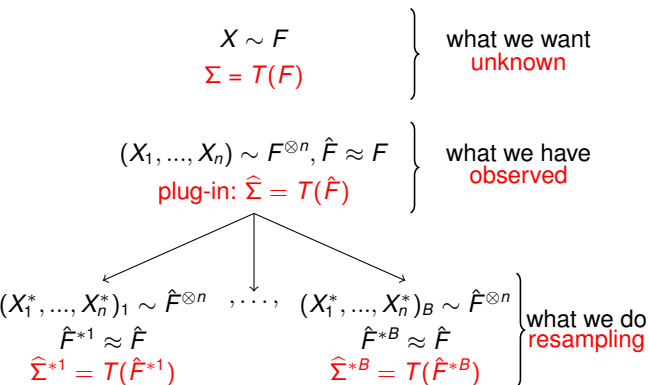
Example:

$$\begin{aligned} (X_1, \dots, X_n) &= (1, 2, 3, 4, 5) \\ \text{(or } (X_1, \dots, X_n) &= n \text{ evaluations)} \end{aligned}$$

Bootstrap: general principle



Bootstrap: general principle



Example:

$(X_1, \dots, X_n) = (1, 2, 3, 4, 5)$
 (or $(X_1, \dots, X_n) = n$ evaluations)

$(X_1^*, \dots, X_n^*)_1 = (1, 3, 4, 3, 5)$
 $(X_1^*, \dots, X_n^*)_2 = (4, 1, 2, 3, 4)$

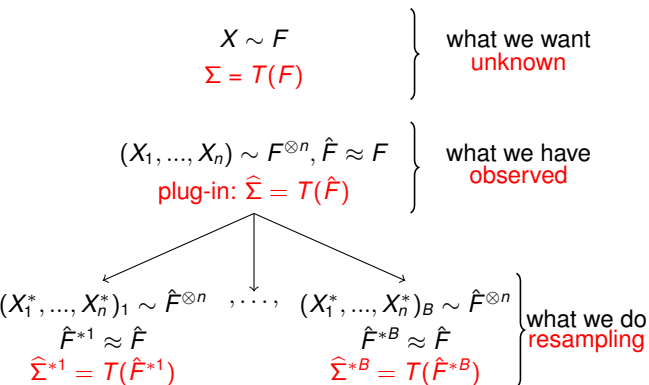
$\vdots$

$(X_1^*, \dots, X_n^*)_B = (5, 1, 3, 4, 2)$

B resamplings allow statistics on the distribution of $\hat{\Sigma}$:

$$\widehat{E}(\hat{\Sigma}), \widehat{\text{Var}}(\hat{\Sigma}), \widehat{\text{Bias}}(\hat{\Sigma}), \dots$$

Bootstrap: general principle



Example:

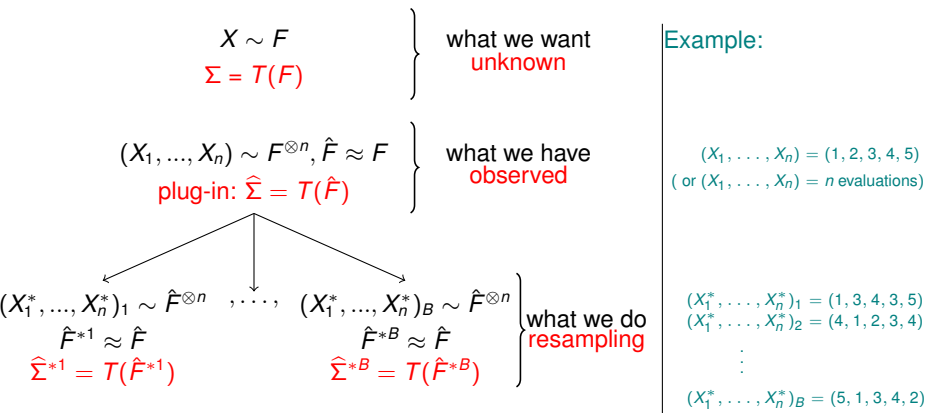
$(X_1, \dots, X_n) = (1, 2, 3, 4, 5)$
 (or $(X_1, \dots, X_n) = n$ evaluations)

$(X_1^*, \dots, X_n^*)_1 = (1, 3, 4, 3, 5)$
 $(X_1^*, \dots, X_n^*)_2 = (4, 1, 2, 3, 4)$
 $\vdots$
 $(X_1^*, \dots, X_n^*)_B = (5, 1, 3, 4, 2)$

B resamplings allow statistics on the distribution of $\hat{\Sigma}$:

$$d_E(\hat{\Sigma}, \Sigma) \approx \widehat{d_E}(\hat{\Sigma}, \Sigma) = \frac{1}{B} \sum_{b=1}^B d_E(\hat{\Sigma}^{*b}, \hat{\Sigma}) \quad \text{and} \quad d_K(\hat{\Sigma}, \Sigma) \approx \widehat{d_K}(\hat{\Sigma}, \Sigma) = \frac{1}{B} \sum_{b=1}^B d_K(\hat{\Sigma}^{*b}, \hat{\Sigma})$$

Bootstrap: general principle



B resamplings allow statistics on the distribution of $\hat{\Sigma}$:

$$d_E(\hat{\Sigma}, \Sigma) \approx \widehat{d_E}(\hat{\Sigma}, \Sigma) = \frac{1}{B} \sum_{b=1}^B d_E(\hat{\Sigma}^{*b}, \hat{\Sigma}) \quad \text{and} \quad d_K(\hat{\Sigma}, \Sigma) \approx \widehat{d_K}(\hat{\Sigma}, \Sigma) = \frac{1}{B} \sum_{b=1}^B d_K(\hat{\Sigma}^{*b}, \hat{\Sigma})$$

$\widehat{d_E}(\hat{\Sigma}, \Sigma)$ and $\widehat{d_K}(\hat{\Sigma}, \Sigma)$ are independant of the method used for the generation of $\hat{\Sigma}$

Bootstrap: theoretical justification

Bickel and Freedman 1981

if $\sqrt{n}(T(\hat{F}) - T(F)) \xrightarrow{\mathcal{L}} P$ then $\sqrt{n}(T(\hat{F}^*) - T(\hat{F})) \xrightarrow{\mathcal{L}} P$

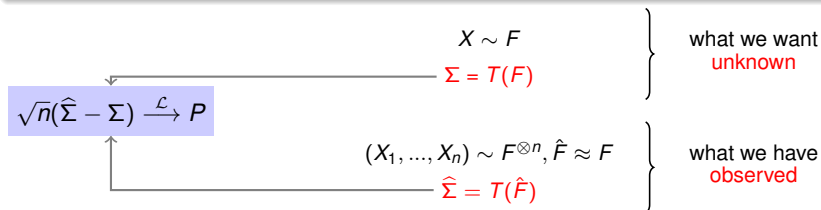
Bickel and Freedman, *Some asymptotic theory for the bootstrap*, the Annals of Statistics, 1981.

Efron, *The Jackknife, the Bootstrap and Other Resampling Plans*, Society for Industrial and Applied Mathematics, 1981.

Bootstrap: theoretical justification

Bickel and Freedman 1981

if $\sqrt{n}(T(\hat{F}) - T(F)) \xrightarrow{\mathcal{L}} P$ then $\sqrt{n}(T(\hat{F}^*) - T(\hat{F})) \xrightarrow{\mathcal{L}} P$



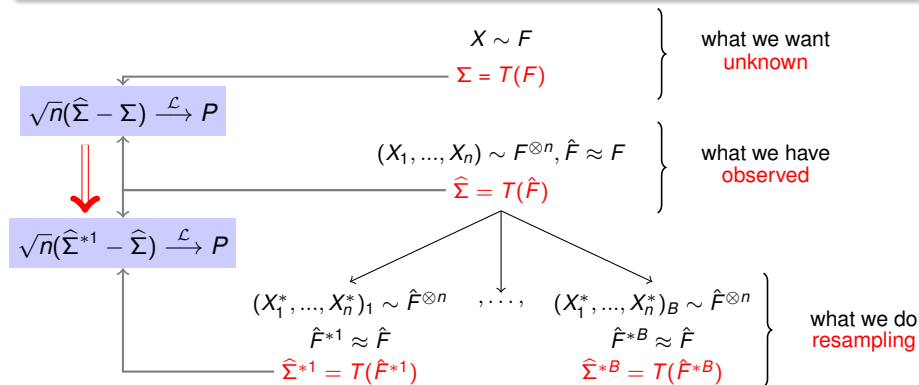
Bickel and Freedman, *Some asymptotic theory for the bootstrap*, the Annals of Statistics, 1981.

Efron, *The Jackknife, the Bootstrap and Other Resampling Plans*, Society for Industrial and Applied Mathematics, 1981.

Bootstrap: theoretical justification

Bickel and Freedman 1981

if $\sqrt{n}(T(\hat{F}) - T(F)) \xrightarrow{\mathcal{L}} P$ then $\sqrt{n}(T(\hat{F}^*) - T(\hat{F})) \xrightarrow{\mathcal{L}} P$



Bickel and Freedman, *Some asymptotic theory for the bootstrap*, the Annals of Statistics, 1981.

Efron, *The Jackknife, the Bootstrap and Other Resampling Plans*, Society for Industrial and Applied Mathematics, 1981.

Bootstrap: theoretical justification

Beran and Strivastava 1985

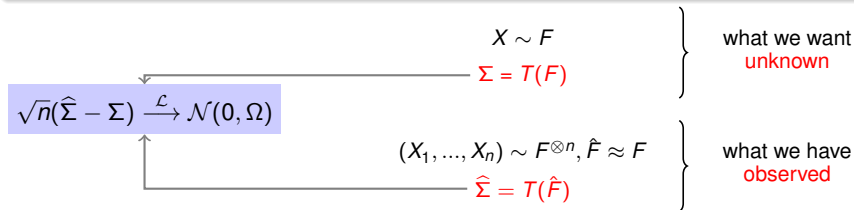
with $\hat{F}$ the empirical CDF, $\sqrt{n}(\hat{\Sigma} - \Sigma) \xrightarrow{\mathcal{L}} \mathcal{N}(0, \Omega)$ and $\sqrt{n}(\hat{\Sigma}^* - \hat{\Sigma}) \xrightarrow{\mathcal{L}} \mathcal{N}(0, \Omega)$

Beran and Strivastava, *Bootstrap tests and confidence region for functions of covariance matrix*, The Annals of Statistics, 1985.

Bootstrap: theoretical justification

Beran and Strivastava 1985

with $\hat{F}$ the empirical CDF, $\sqrt{n}(\hat{\Sigma} - \Sigma) \xrightarrow{\mathcal{L}} \mathcal{N}(0, \Omega)$ and $\sqrt{n}(\hat{\Sigma}^* - \hat{\Sigma}) \xrightarrow{\mathcal{L}} \mathcal{N}(0, \Omega)$

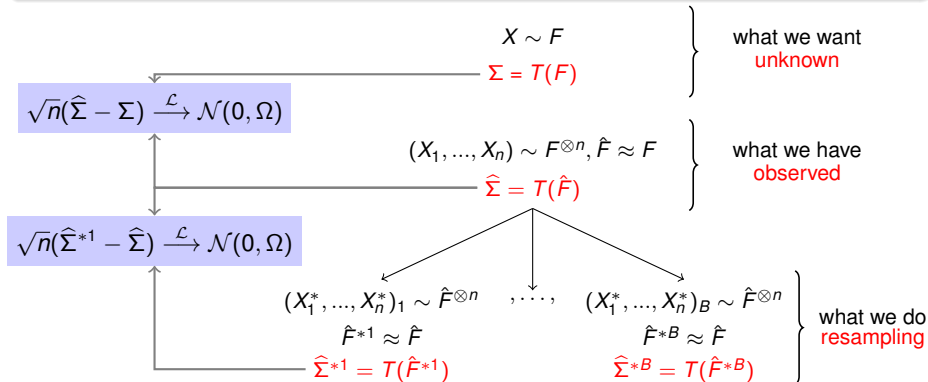


Beran and Strivastava, *Bootstrap tests and confidence region for functions of covariance matrix*, The Annals of Statistics, 1985.

Bootstrap: theoretical justification

Beran and Strivastava 1985

with $\hat{F}$ the empirical CDF, $\sqrt{n}(\hat{\Sigma} - \Sigma) \xrightarrow{\mathcal{L}} \mathcal{N}(0, \Omega)$ and $\sqrt{n}(\hat{\Sigma}^* - \hat{\Sigma}) \xrightarrow{\mathcal{L}} \mathcal{N}(0, \Omega)$

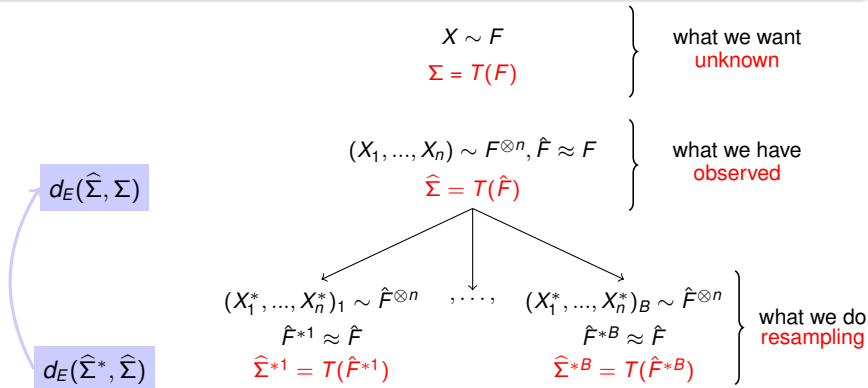


Beran and Strivastava, *Bootstrap tests and confidence region for functions of covariance matrix*, The Annals of Statistics, 1985.

Bootstrap: theoretical justification

Beran and Strivastava 1985

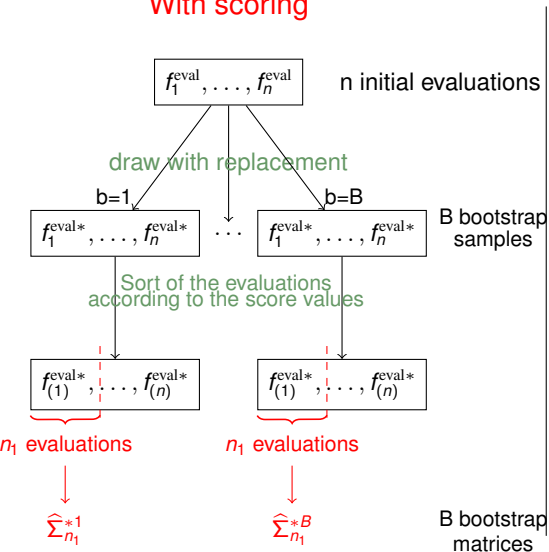
with $\hat{F}$ the empirical CDF, $\sqrt{n}(\hat{\Sigma} - \Sigma) \xrightarrow{\mathcal{L}} \mathcal{N}(0, \Omega)$ and $\sqrt{n}(\hat{\Sigma}^* - \hat{\Sigma}) \xrightarrow{\mathcal{L}} \mathcal{N}(0, \Omega)$



Beran and Strivastava, *Bootstrap tests and confidence region for functions of covariance matrix*, The Annals of Statistics, 1985.

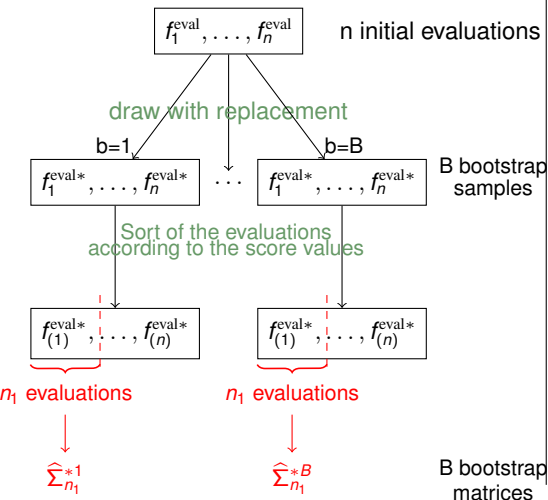
Bootstrap implementation

With scoring

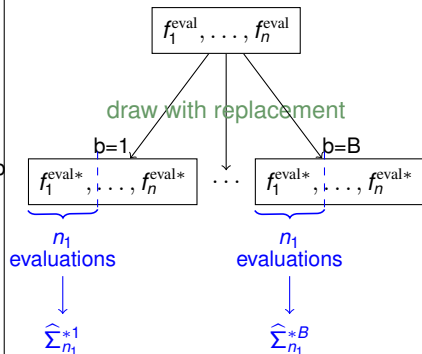


Bootstrap implementation

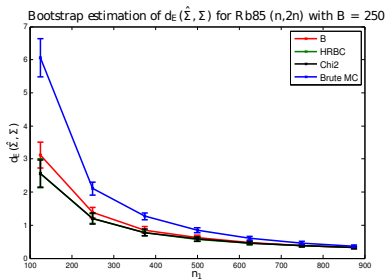
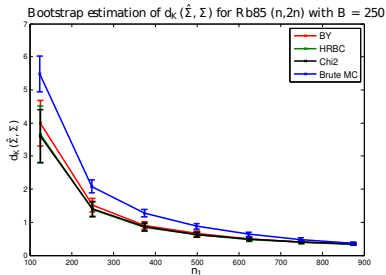
With scoring



With Brute Monte Carlo



Entropy and Kullback-Leibler versus n_1

 $\hat{d}_E$

 $\hat{d}_K$


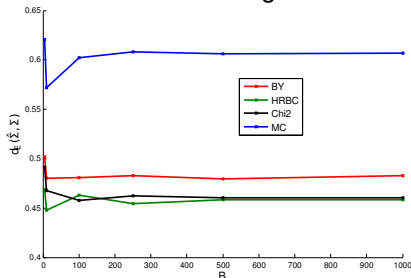
Entropy ratio and Kullback-Leibler ratio versus B

$n_1 = 625$

Local scoring

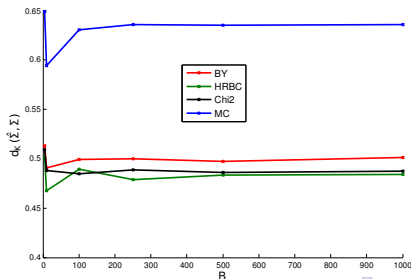
$d_E(\hat{\Sigma}, \Sigma)$

for Rb85 (n,2n)



$d_K(\hat{\Sigma}, \Sigma)$

for Rb85 (n,2n)



En. ratio and K-L ratio versus param. variation

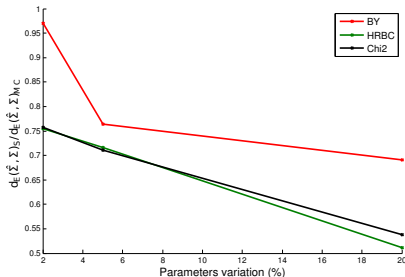
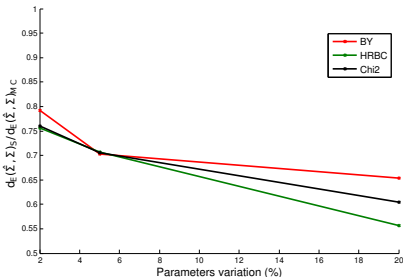
$B = 500,$
 $n_1 = 625$

Local scoring

Global scoring $((n, 2n), (n, \gamma), (n, \alpha))$

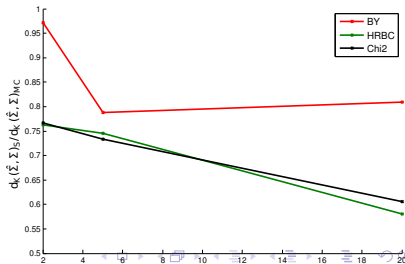
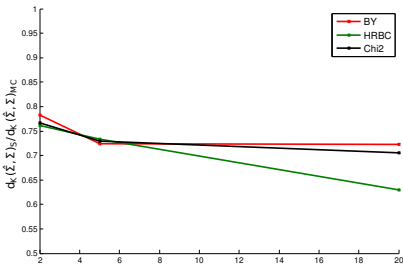
$$\frac{d_E(\hat{\Sigma}, \Sigma)_S}{d_E(\hat{\Sigma}, \Sigma)_{MC}}$$

for Rb85
 $(n, 2n),$



$$\frac{d_K(\hat{\Sigma}, \Sigma)_S}{d_K(\hat{\Sigma}, \Sigma)_{MC}}$$

for Rb85
 $(n, 2n)$



References



A. J. Koning, J. P. Delaroche, *Local and global nucleon optical models from 1 keV to 200 MeV*, Nucl. Phys. A713, 231 (2003).



O. Bersillon, *The computer Code SCAT2000*, Workshop on Applied Nuclear Theory and Nuclear Model Calculations for Nuclear Technology Applications, M.K. Mehta and J.J.Schmidt editors, 1988.



S. Wiener, *The Homogeneous Chaos*, American Journal of Mathematics, (1938).



R. Ghanem and P. Spanos. *Stochastic finite elements : a spectral approach*. Springer, Berlin, 1991.



O. Le Maître, O. Knio, H. Najm and R. Ghanem, *Multi-resolution analysis of Uncertainty propagation schemes using Wiener-type expansions*, Journal of Computational Physics, (2003).



R.H. Cameron, W.T. Martin, *The orthogonal development of nonlinear functionals in series of Fourier-Hermite functionals*, Ann.Math., 48 : 385-392 (1947).



B. Sudret, *Global sensitivity analysis using polynomial chaos expansion*, 2007.



M. B. Chadwick, T. Kawano, P. Talou, E. Bauge, S. Hilaire, P. Dossantos-Uzarralde, P. E. Garrett, *Yttrium ENDF/B-VII Data from Theory and LANSCE/GEANIE Measurements and Covariances Estimated using Bayesian and Monte-Carlo Methods* Nuclear Data Sheets 108 (2007) 2742 - 2755



P. Dossantos-Uzarralde, A. Guittet *A Polynomial Chaos Approach for Nuclear Data Uncertainties Evaluations* Nuclear Data Sheets 109 (2008) 2894 - 2899