

— 8TH INTERNATIONAL CONFERENCE —



SENSITIVITY ANALYSIS OF MODEL OUTPUT

Celebrating the 90th birthday of ILYA M. SOBOL'

PROCEEDINGS

Conference co-organised by the

PHYSICS AND MATHEMATICAL ENGINEERING LABORATORY
FOR ENERGY, ENVIRONMENT AND BUILDING (PIMENT, EA 4518)

&

THE JOINT RESEARCH CENTER (EUROPEAN COMMISSION)

AT THE UNIVERSITY OF REUNION ISLAND, LE TAMPON

November 30th to December 3rd, 2016

Table of contents

Simultaneous estimation of groundwater recharge and hydrodynamic parameters for groundwater flow modelling, Ackerer Philippe [et al.]	1
Sobol' indices and variance reduction diagram estimation from samples used for uncertainty propagation, Benaichouche Abed	3
Getting better insights in the influence of uncertainties in seismic risk. Application to L'Aquila earthquake (2009)., Benaichouche Abed	6
Dynamic coupling between μ CHP systems and buildings: sensibility analysis of the time resolution of the electrical demand data and of μ CHP modeling typology, Bouvenot Jean-Baptiste [et al.]	9
Quantile-oriented sensitivity indices, Browne Thomas	11
Investigation of Modern Methods of Probabilistic Sensitivity Analysis of Final Repository Performance Assessment Models, Becker Dirk-Alexander	13
Optimising Composite Indicators with Sensitivity Analysis, Becker William [et al.]	15
A Copula-based Approach to Sensitivity to Correlations in Structural Reliability Problems, Benoumechiara Nazih [et al.]	17
Sensitivity analysis under different distributions using the same simulation, Bolado-Lavin Ricardo [et al.]	19
Predicted Sensitivity for establishing well-posedness conditions in stochastic inversion problems, Bousquet Nicolas [et al.]	21
Sensitivity analysis and the calibration problem of a biodynamic model for Indian Ocean tuna growth, Bousquet Nicolas [et al.]	23
Link between the sensitivity indices at different scales, Bruchou Claude	25
Using gaussian process metamodels for sensitivity analysis of an individual-based model of a pig fattening unit, Cadero Alice [et al.]	27
Sensitivity analysis and calibration of a numerical code for the prediction of power from a photovoltaic plant, Carmassi Mathieu [et al.]	30

A new Bayesian approach for statistical calibration of computer models, Delay Frederick [et al.]	32
Adaptive numerical designs for the calibration of computer codes, Damblin Guillaume . .	34
Global Sensitivity analysis in dynamic MFA, Dzubur Nada	36
Development of a High Performance Capabilities for Supporting Spatially-Explicit Uncertainty- and Sensitivity Analysis in Multi-Criteria Decision Making, Erlacher Christoph [et al.] . .	38
Sensitivity and uncertainty analyses in climate research: Tool development and application for an atmosphere model, Flechsig Michael [et al.]	40
Using the sensitivity analysis to optimize passive cooling solutions in the urban tropical environment, Fouquier Aurelie [et al.]	42
EnergyPlus Laboratory for Sensitivity and Uncertainty Analysis in Building Energy Modeling : The EPLab Software, Goffart Jeanne [et al.]	44
Dynamic sensitivity analysis method based on Gramian, Hamza Sabra [et al.]	46
Sobol' sensitivity analysis for stochastic numerical codes, Iooss Bertrand [et al.]	48
Numerical stability of Sobol' indices estimation formula, Iooss Bertrand [et al.]	50
Sobol' indices for problem defined in non-rectangular domains, Kucherenko Sergei [et al.]	52
Global sensitivity analysis of non-domestic analysis buildings thermal behaviour, Kucherenko Sergei [et al.]	54
A minimum variance unbiased (generalized) estimator of total sensitivity indices: an illustration to a flood risk model, Lamboni Matieyendou	56
Sensitivity analysis and metamodeling methods for designing buffer strips to protect water from pesticide transfers., Lauvernet Claire [et al.]	58
The use of computed assisted semen analysis (CASA) as a method for environmental and toxicological risk assessment – The use of different chambers as sensitivity factor, Massanyi Peter [et al.]	61
The role of Rosenblatt transformation in global sensitivity analysis of models with dependent inputs, Mara Thierry	63
PC Expansion for Global Sensitivity Analysis of non-smooth funtionals of uncertain Stochastic Differential Equation solutions, Navarro Jimenez Maria [et al.]	65
Statistical emulation as a tool for analysing complex multiscale stochastic biological model outputs, Oyebamiji Oluwole	67
Identification of influential parameters in building energy simulation and life cycle assessment, Pannier Marie-Lise [et al.]	69

Sensitivity analysis as essential tool to gain insight into potential hydrological change due to coal development in Australia., Peeters Luk	71
Global sensitivity analysis with distance correlation and energy statistics, Plischke Elmar [et al.]	73
Combining switching factors and filtering operators in GSA to analyze models with climatic inputs, Roux Sebastien [et al.]	75
RepoSTAR –New Framework for Statistic Runs for Uncertainty and Sensitivity Analysis of a Radioactive Waste Repository Model, Reiche Tatiana [et al.]	77
Sensitivity Analysis for Energy Performance Contracting in new buildings, Rivalin Lisa [et al.]	79
Investigating Scale Effect of Watershed Delineations on Local Multi-Criteria Method for Land Use Evaluation, Salap-Ayca Seda [et al.]	81
PCE-based Sobol’ indices for probability-boxes, Schoebi Roland [et al.]	83
Bayesian Sparse Polynomial Chaos Expansion, Shao Qian [et al.]	85
A guideline for sensitivity analysis of repository models, Spiessl Sabine M. [et al.]	87
Perturbed-Law based sensitivity Indices for sensitivity analysis in structural reliability, Sueur Roman	89
Variance-based sensitivity analysis for inputs over non-rectangular domains, Tarantola Stefano [et al.]	91
Confidence intervals for Sobol’ indices, Touati Taieb	93
Global sensitivity analysis by HDMR combining with improved GMDH algorithm, Wang Lu [et al.]	95
A new method of network clustering based on second-order sensitivity index, Wu Qiongli [et al.]	97
Model Interpretation via Multivariate Padé-Approximant in Low-Rank Format, Zivanovic Rastko	99
Fixed and sequential designs for optimisation of fan shapes, Azais Jean-Marc [et al.]	101
New Fréchet features for random distributions and associated sensitivity indices, Fort Jean Claude	102
Comparison of Latin Hypercube and Quasi Monte Carlo Sampling Techniques, Saltelli Andrea [et al.]	105
Global sensitivity analysis and Bayesian parameter inference for transport in a dual flowing continuum, Younes Anis [et al.]	108

Real-time building design space exploration using two-sample Kolmogorov Smirnov tests
to rank inputs according to multiple outputs, Ostergaard Torben [et al.] 110

List of participants 112

Author Index 112

Simultaneous estimation of groundwater recharge and hydrodynamic parameters for groundwater flow modelling

F.Z. Hassane Mamadou Maina^{1,2}, O. Bilstein², and P. Ackerer¹

¹Laboratoire Hydrologie et Geochimie de Strasbourg, University of Strasbourg/EOST, CNRS, 1 rue Blessig
67084 Strasbourg, France

²EA-Laboratoire de Modélisation des Transferts dans l'Environnement, Bât. 225, F-13108 Saint Paul lez
Durance cedex, France

Abstract

Groundwater water resources management is a major issue because it is often the only available water resources for many countries. Moreover, the estimation of the recharge of the groundwater is still a challenge and is a key value for accurate exploitation of the aquifer. Therefore, groundwater models have become a very usual tool for groundwater management. Aquifers are known to be very heterogeneous and groundwater recharge to be variable in space, depending mainly on the soil and vegetation covers, and in time since it is a combination of precipitation, evaporation and transpiration through the vegetation. Groundwater recharge cannot be measured. Its value varies from zero to precipitation. Two key factors are unknown: the actual evapotranspiration which depends on many factors like vegetation and meteorological data and the water flow conditions in the unsaturated zone located above the groundwater. Groundwater survey consists in measuring water piezometric levels in different wells. Physical parameters (hydraulic conductivity, storage coefficient), aquifer geometry and boundary conditions are very poorly known for economic reasons: the parameters estimation requires the drilling and monitoring of numerous wells. Therefore, model calibration cannot be avoided. Groundwater models are based on two equations: - Mass conservation written in terms of piezometric head assuming constant water density:

$$S \frac{\partial h}{\partial t} + \nabla \mathbf{q} = r$$

- and Darcy's law for energy conservation:

$$\mathbf{q} = -\mathbf{K} \nabla h$$

where h is the piezometric head, S the storage coefficient, $\mathbf{K}$ the hydraulic conductivity tensor, $\mathbf{q}$ the water flux and r sink/source terms including groundwater recharge.

Model calibration consists in fitting the measured piezometric heads by estimating the ad hoc parameters (storage term and hydraulic conductivity) and sink/source terms. Boundary conditions are rarely calibrated. It is traditionally recommended to avoid simultaneous calibration of groundwater recharge and flow parameters because of correlation between recharge and hydraulic conductivity. From a physical point of view, little recharge associated with low hydraulic conductivity can provide very similar piezometric values than high recharge and high hydraulic conductivity.

If this correlation is true under steady state conditions, we assume that this correlation is much weaker under transient conditions because recharge varies in time and the parameters do not. Moreover, the

recharge is negligible during summer time for many climatic conditions due to reduced precipitation, increased evaporation and transpiration by vegetation cover.

We analyze our hypothesis through global sensitivity analysis (GSA) in conjunction with the polynomial chaos expansion (PCE) methodology. We perform GSA by calculating the Sobol indices, which provide a variance-based importance measure of the effects of uncertain parameters (storage and hydraulic conductivity) and sink/source term on the piezometric heads computed by the flow model. The choice of PCE has the following two benefits: (i) it provides the global sensitivity indices in a straightforward manner, and (ii) PCE can serve as a surrogate model for the calibration of parameters. The coefficients of the PCE are computed by probabilistic collocation. We perform the GSA on real conditions coming from an already built groundwater model dedicated to a subdomain of the Upper-Rhine aquifer (geometry, boundary conditions, climatic data, measured piezometric heads over time in several wells).

GSA shows that the simultaneous calibration of recharge and flow parameters is possible if the calibration is performed over at least one year. It provides also the valuable information of the sensitivity versus time, depending on the aquifer inertia and climatic conditions. Typically, piezometric heads are sensitive to flow parameters when recharge is negligible (summer time) and are more sensitive to recharge in winter time.

Our work shows also that the GSA method in conjunction with the PCE technique can provide practical guidance for groundwater resources survey.

Sobol'indices and variance reduction diagram estimation from samples used for uncertainty propagation

A. Benaïchouche¹ and J. Rohmer¹

¹BRGM, The French Geological Survey, France

Abstract

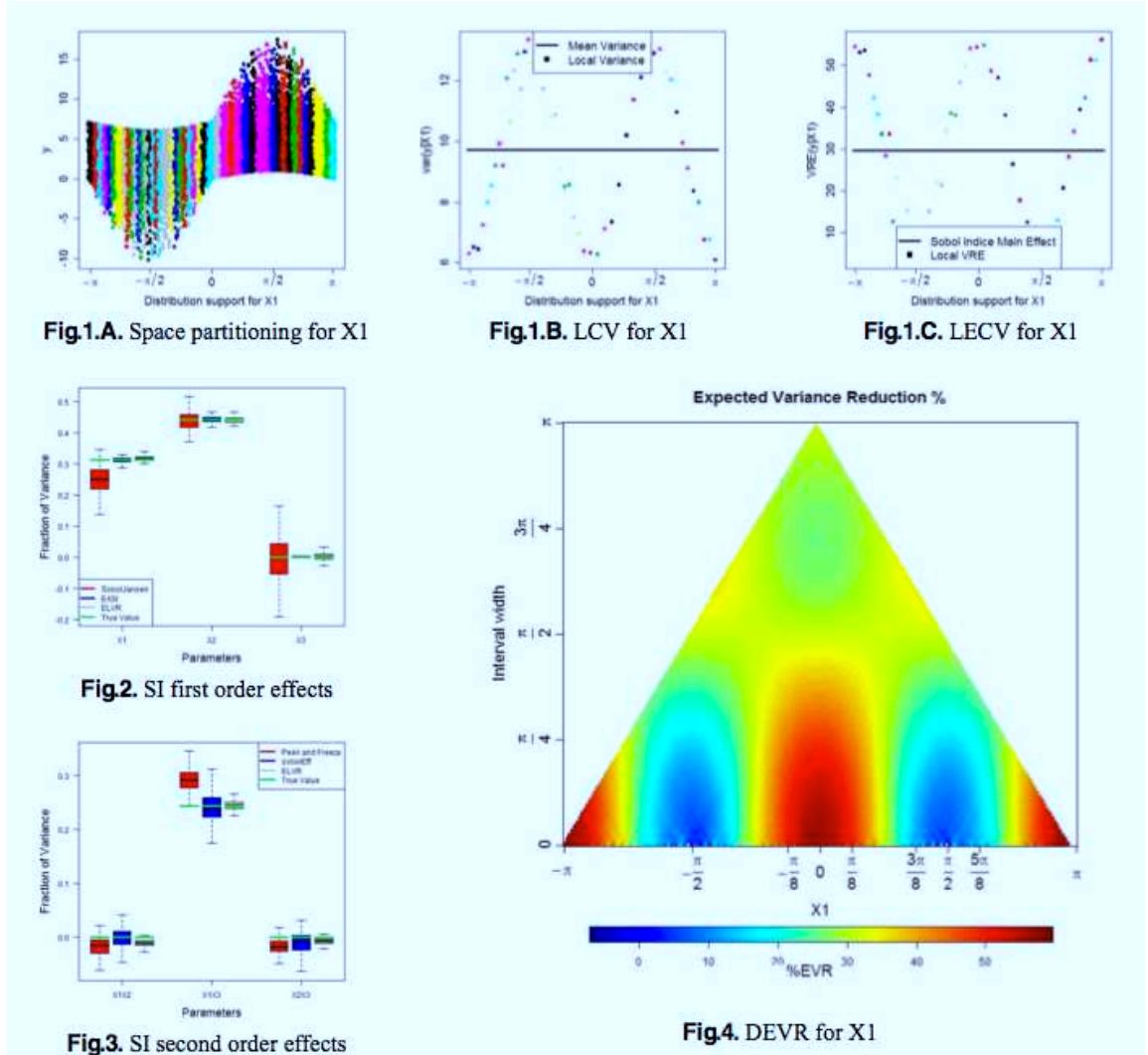
In this paper we present an efficient algorithm to estimate first order and second order Sobol'indices (SI) and its relationships with the inputs parameters variation range width. In the variance-based global sensitivity measure, SI are quantities defined by normalizing parts of variance in ANOVA decomposition (Sobol', 1993). They are estimated from the ratio between the variance of the conditional expectation of the output given the input and the unconditional variance of the output (1). Many techniques have been proposed to estimate these indices. A recent review of these methods can be found in (Borgonovo & Plischke, 2015). Among others Monte-Carlo-based algorithm (Sobol' 1993, Saltelli et al., 1999, Jansen 1999, Monod et al., 2006, etc.) require a special random sampling scheme i.e. they cannot directly use samples used for the uncertainty propagation. Building on a similar idea than Plischke (Plischke, 2010) we propose a methodology that allows the computation of SI from a set of given data. We propose to start from the local conditional variance to derive the global sensitivity indicator by estimating the variance of the conditional expectation (2). The local information on sensitivity is summarized under the form of a Diagram of Expected Variance Reduction (DEV), which relates the local reduction in uncertainty with the domain of variation of the considered input parameter with reduced width.

SI are given by (1). The variance of the conditional expectation can be expressed as the expectation of conditional variance (2). In order to estimate the variance of the conditional expectation $V[E[(Y|x_i)]]$, we first estimate the expectation of the conditional variance $E[V[(Y|x_i)]]$ by following these steps: (i) Partition of the input parameters space into clusters (K -means algorithm with fixed number of samples for example), in which the variation of x_i is supposed weak (see fig. 1.A). (ii) Computation of the local conditional variance $V[(Y|x_i)]$ (LCV) for each cluster (see fig. 1.B). (iii) Estimation of the expectation of local conditional variance $E[V[(Y|x_i)]]$ (ELCV) with (2) (fig. 1.C). (iv) SI are obtained from the average of the expectation of local conditional variance. Higher order effects can be estimated following the same scheme.

$$S_i = \frac{V[E[(Y|x_i)]]}{V[Y]} \quad (1)$$

$$V[E[(Y|x_i)]] = V[Y] - E[V[(Y|x_i)]] \quad (2)$$

The described methodology allows also the computation of the relationship between SI and the inputs parameters variation range width (DEV). In other words, this answers the question of what would be SI values if the initial variation range was reduced or true value of parameters was known: this can be achieved without new simulations run contrary to classical methods.



The proposed algorithm is used to estimate the first order and second order SI and the DEVR on the Ishigami function (3). This widely used test function exhibits strong nonlinearity and non-monotonicity (Sobol' & Levitan 1999) which make it challenging.

$$y = f(x_1, x_2, x_3) = \sin(x_1) + 7 \sin^2(x_2) + 0.1x_3^4 \sin(x_1), \text{ where } x_i \sim U(-\pi, \pi) \quad (3)$$

We tested the developed algorithm termed ELVR against Jansen's algorithm (Jansen, 1999) and EASI (Plischke, 2010) for the first effects (5,000 model run) and against peek and freeze algorithm (Fruth & al 2014) and Monod's formulation (Monod et al. 2006) for the second order effects (10,000 model run). The results are reported in fig. 2 and fig. 3: these show a very good convergence for the first order SI index with less than 10,000 simulations as well as for the second-order indices especially for indices of high values.

From the DEVR (fig. 4) many information on local sensitivity can be extracted. Among others:

- In fig. 2, $S_1 = 0.31$, which means that if x_1 is fixed; a reduction of 31% on the total variance is expected. But in the DEVR (fig. 4) we see that depending on the chosen value the reduction can vary between 55% ($x_1^* = -\pi, 0, \pi$) and 0% ($x_1^* = -\pi/2, \pi/2$). But on average it's equal to 31%.

- The reduction of variation range of x_1 from $[-\pi, \pi]$ to $[-\pi/8, \pi/8]$ reduces the total variance of output by more than 50%. On the other hand, a reduction from $[-\pi, \pi]$ to $[3\pi/8, 5\pi/8]$ do not change $V[Y]$.

Acknowledgement: This work has been funded by the French National Research Agency (ANR) under action 2 of the convention # 14 CARN 013-01 with Institut Carnot BRGM.

References

- Cukier, R. I., Levine, H. B., and Shuler, K. E. (1978). Nonlinear sensitivity analysis of multiparameter model systems. *Journal of Computational Physics*, 26:1-42.
- J. Fruth, O. Roustant, S. Kuhnt, 2014, Total interaction index: A variance-based sensitivity index for second-order interaction screening, *J. Stat. Plan. Inference*, 147, 212-223.
- M.J.W. Jansen, 1999, Analysis of variance designs for model output, *Computer Physics Communication*, 117, 35-43.
- Monod, H., Naud, C., Makowski, D. (2006), Uncertainty and sensitivity analysis for crop models in *Working with Dynamic Crop Models: Evaluation, Analysis, Parameterization, and Applications*, Elsevier.
- E. Plischke, 2010, An effective algorithm for computing global sensitivity indices (EASI). *Reliability Engineering & System Safety* 95, 354-360.
- Saltelli, A., Tarantola, S., and Chan, K. P. S. (1999). A quantitative model-independent method for global sensitivity analysis of model output. *Technometrics*, 41:39-56.
- Sobol', I. M. (1993). Sensitivity analysis for nonlinear mathematical models. *Mathematical Modeling and Computational Experiment*, 1:407-414.
- Sobol', I. M. & Levitan, Y. L. (1999). On the use of variance reducing multipliers in Monte Carlo computations of a global sensitivity index. *Computer Physics Communications*, 117(1), 52-61.
- Tarantola, S., Gatelli, D., and Mara, T. A. (2006). Random balance designs for the estimation of first-order global sensitivity indices. *Reliability Engineering and System Safety*, 91:717-727.
- E. Borgonovo, E. Plischke, Sensitivity analysis: A review of recent advances. *European Journal of Operational Research*. In press (2015).

Getting better insights in the influence of uncertainties in seismic risk. Application to L'Aquila earthquake (2009)

A. Benaïchouche¹, J. Rohmer¹, D. Monfort Climent¹, and Christian Bellier¹

¹BRGM, The French Geological Survey, France

Abstract

Over recent years, numbers of tools for seismic risk analysis have been developed to evaluate casualties and losses induced by earthquakes. A recent overview of available models can be found in (Molina et al. 2010). These predictions software require a large number of quantitative parameters (parametric uncertainty) but also model structures (model uncertainty). The choice of the appropriate model and the determination of exact values of these parameters remain very difficult. Therefore, sensitivity analysis and uncertainty quantification have to be performed for these kinds of studies. In this context, the aim of this article is to present a methodology for getting better insight in the role played by the different uncertainty sources (parametric and model) based on a variance-based global sensitivity analysis. Contrary to Rohmer et al. (2014), we used a less greedy estimation algorithm (Benaïchouche & Rohmer 2016), which both allows providing global sensitivity measures, but also information on local sensitivity. The application case is the risk analysis performed for Aquila (Italy, 2009) earthquake.

The seismic damage assessment models evaluate earthquake-related risk, casualties, and losses through the convolution of two independent modules: seismic aggression and vulnerability. For seismic aggression estimation three steps are required: (i) Regional hazard estimation by calculating the ground shaking (Peak Ground Acceleration on bedrock PGA) induced an earthquake defined by its epicenter position (XY), depth (Z), magnitude (Mw) using ground motion prediction equations (GMPE) associated to a statistical parameter (Sigma). (ii) PGA at local scale, which accounts for lithological effects via an amplification factor (LITH). (iii) The PGA value is converted in macro-seismic intensity using conversion laws (GMICE). The obtained intensity is then convoluted with vulnerability module for damage estimation. In this work, we focus on the sensitivity analysis related to the estimated intensity output. All our simulations were performed with Armagedom software (Sedan et al, 2013).

To assess the sensitivity analysis of the intensity output we developed a new technique for Sobol' indices estimation, where the variance of the conditional expectation $V[E[(Y|x_i)]]$ is calculated from the expectation of the local conditional variance $E[V[(Y|x_i)]]$ by partitioning the input parameters space into clusters. This algorithm (ELVR) is presented in details in (Benaïchouche & Rohmer, 2016). This method allows the estimation of the first order and second order Sobol' indices and its relationships with the inputs parameters (respectively: model) variation range width (respectively: model choice). Here, the obtained results with the developed technique will be presented and compared with those obtained with Saltelli's algorithm (Saltelli, 2002) for the first order Sobol' indices estimation.

L'Aquila is a moderate-sized city ($\sim 73,000$ inhabitants) located in Central Italy (Fig.1), 90km in the south-west of Roma. On 6 April 2009, a 6.3 earthquake magnitude hit the region causing severe damages: 308 people have died, 20% of the housing was heavily damaged and more than 40,000 people were left homeless (Tertuliani, 2010). The impact on religious and monumental heritage was

also disastrous. The estimated intensity in the L'Aquila center is around 8.5 (EMS scale). In this context, our objective is the evaluation of uncertainties propagation (parameter and model) related to the average estimated intensity in the L'Aquila. The uncertainties sources for this application are summarized in table 1. The estimated values were obtained from Douglas et al. 2015. More details are provided in Douglas et al. (2015) and references therein.

The quantity of interest is the intensity (average over the whole area of the L'Aquila Historical city) for the Sobol' indices estimation with 7 random inputs (Table.1). Samples were generated using Sobol' quasi-random sequences technique. The performed simulations give an average intensity of 7.5 and a total variance of 1 (Fig.1 show the results of an arbitrary simulation). The results are given in Fig. 2 (left). Fig.2 (left) shows that the developed method requires a fewer number of model evaluations (1K simulations) to obtain an excellent convergence compared to those obtained with Saltelli algorithm (200K simulations).

Analysis of results of the first order effects (Fig.2-left) shows that the model is additive (sum of seven main effects is 95.6%). This means that all input parameters have a negligible interaction effects between them. This result is confirmed in Fig.2 (middle), which shows that all second order Sobol' indices effects are less than 4% (10K model evaluations). Therefore, the analysis can be restricted to the first order effects. We observe (Fig.2 - left) that the most influential input corresponds to the choice of the GMICE model with a main effect of 58% followed by the GMPE model with a main effect of 14%. Uncertainties on Magnitude (Mw) and Sigma represent a main effect of 8%. Finally, the position, depth and lithological site effects (denoted XY, Z and LITH) have a little influence (less than 3%). These results show that the appropriate choice of GMICE and GMPE models should be prioritized in future investigations.

A local analysis of the first order effects of the GMICE model shows (Fig.2 - right) that the GMICE equations 1,2,3 have respectively an influence of 22%,73%,80%. This means that if additional information were available and the GMICE is fixed to 3, the total variance on the average intensity could be reduced up to 80%. On the other hand, if it was fixed to 1 a reduction of 22% could be expected. The average reduction on total variance is given by the mean of these three values (58.3%) which correspond to first order Sobol' index obtained in Fig.2 (left).

Acknowledgement: This work has been funded by the French National Research Agency (ANR) under action 2 of the convention # 14 CARN 013-01 with Institut Carnot BRGM.

References

- A. Benaïchouche et J. Rohmer, 2016. Sobol' indices and variance reduction diagram estimation from samples used for uncertainty propagation. SAMO 2016.
- J. Douglas, D. Monfort Climent, C. Negulescu, A. Roulleé, O Sedan, 2015. Limits on the potential accuracy of earthquake risk evaluations using the L'Aquila (Italy) earthquake as an example *Annals of Geophysics* 58 (2), 0214.
- Rohmer, J. Douglas, J. Bertil, D. Monfort Climent, & O. Sedan, 2014. Weighing the importance of model uncertainty against parameter uncertainty in earthquake loss assessments. *Soil Dynamics and Earthquake Engineering*, 58, 1–9.
- A. Saltelli, 2002. Making best use of model evaluations to compute sensitivity indices, *Computer Physics Communication*, 145, 580-597.
- O. Sedan, C. Negulescu, M. Terrier, A. Roulle?, T. Winter, D. Bertil, 2013. Armagedon a Tool for Seismic Risk Assessment Illustrated with Applications, *Journal of Earthquake Engineering*, 253-281, doi:10.1080/13632469.2012.726604.
- A. Tertulliani, L. Arcoraci, M. Berardi, F. Bernardini, R. Camassi, C. Castellano, S. Del Mese, E. Ercolani, L. Graziani, I. Leschiutta, A. Rossi and M. Vecchi, 2010. An application of EMS98 in a medium-sized city: The case of L'Aquila (Central Italy) after the April 6, 2009 Mw 6.3 earthquake, *B. Earthq. Eng.*, 9 (1), 67-80; doi:10.1007/s10518-010-9188-4

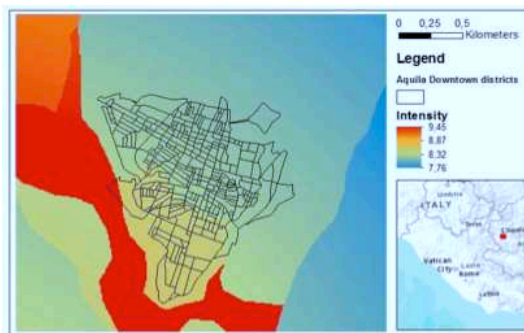


Fig.1. Estimated macro-seismic intensities (IEMS) in L'Aquila Downtown with Armagedom software (Sedan et al., 2013).

Uncertainty source	Representation
XY (Parameter)	Uniform distribution with $\pm 5\text{km}$
Z (Parameter)	Uniform distribution with $\pm 5\text{km}$
Mw (Parameter)	Uniform distribution between 5.5 et 6.5
Sigma (Parameter)	Uniform distribution between -0.5 et 0.5
GMPE (Model)	Discrete variable with variable uniformly taken from $\{1,2,3,4,5,6,7\}$
GMICE (Model)	Discrete variable with variable uniformly taken from $\{1,2,3\}$
LITH (Parameter)	Uniform distribution with 20% from estimated value

Table1. Source of uncertainty description and assumption for representing them in L'Aquila application (XY: Earthquake position, Z: Earthquake depth, Mw: Earthquake magnitude, Sigma: Statistical parameter for GMPE, GMPE: Selection of the GMPE, GMICE: Selection of the GMICE and LITH: Amplification factor of lithological effects).

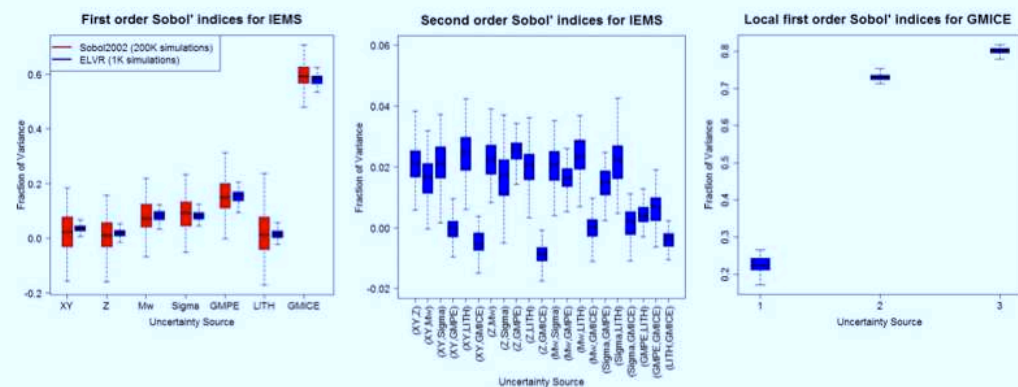


Fig. 2. Boxplot summarizing the influence on sensitivity results for the first order (left), second order (middle) Sobol' indices and the local first order Sobol indices for the GMICE (right) of intensity (IEMS) output. Bottom and top of the boxplot are the 25th and 75th percentile, the ban near the middle of the box is the median and the ends of the whiskers are the minimum and maximum. To obtain confidence intervals (as boxplots) for each estimated indices we perform a bootstrap with 100 replicates.

Dynamic coupling between μ CHP systems and buildings: sensibility analysis of the time resolution of the electrical demand data and of μ CHP modeling typology

J.B. Bouvenot¹, M. Siroux¹, and B. Latour¹

¹ICube UMR 7357, INSA de Strasbourg, 24 Boulevard de la Victoire, 67000 Strasbourg, France

Abstract

The micro combined heat and power (μ CHP) or micro cogeneration is a technology which produces simultaneously decentralized thermal and electrical (or mechanic) energy at low power (electrical power ≤ 50 kWel). This technology recovers the “fatal heat” losses considered as “heat waste” produced in thermodynamics or thermochemical cycles for mechanic energy production. This heat can be used to cover buildings heating and domestic hot water (DHW) needs. The μ CHP matches the two goals of energetic system efficiency and greenhouse gas emission reduction by converting more efficiently the primary energy in final energy [1]. Besides, the integration of these low thermal and electrical power systems within the energy consumption places lets to self-consume the produced energy, to relieve the grid mainly during peak demand hours and to avoid grid losses.

A wood pellet steam engine and a gas (or biogas) Stirling engine μ CHP devices have been tested at the laboratory of INSA Strasbourg in order to characterize their performances in steady and unsteady states. Two realistic and dynamic models based on these experimental investigations have been developed in previous works [2, 3] in order to predict their energy performances and their pollutant emissions. These models have been implemented in the TRNSYS’s numerical environment where an optimization platform has been implemented. Thermal and electrical energy storage systems and energy management controller have been implemented in this platform which is used to optimize the coupling between buildings and this kind of innovative devices by considering energetic, economic and environmental criteria. [4]. Dynamic thermal simulations (DTS) only computes dynamic heating loads but the other most crucial parameters of the platform are the DHW load profiles and mainly the electrical load profiles in buildings which needs to be realistic, variable, suitable to the French context and with a low time step. Existing data basis are weakly suited to our platform because of their lack of precision (more than 5 min time step), their lack of information (no information about the load profiles for each electrical appliance) or their non-relevance in the French context.

Stochastic and high resolution electrical demand and DHW demand generators have been created and are well adapted to the French context by using a “bottom-up” method aggregating the electrical load of each electrical appliance or specific DHW draw-off by a stochastic way [5].

Here we propose two kinds of sensitivity analysis:

The first one deals with the time resolution of the electrical needs. This time resolution appears as crucial to catch the real variation of this quantity which is very variable and volatile. This “electrical” time step is not suitable with the “thermal” time step of thermal building energy simulations which involves usually hourly or semi hourly time steps. A big time step tends to smooth the peak of the electrical demand and artificially increase the self-consumption of the produced electricity. Besides,

we propose a numerical repeatability analysis to show the dispersion about self-consumption ratios linked to different electrical demand data generated by the stochastic algorithm. The aim is to show the impact of the simulation time step and the reliability of the results linked to a data file. This study will help to obtain a reliable and realistic final result.

The second analysis deals with the modeling typology. Precise data driven models have been developed by taking into account transient behavior and boundary conditions influence (cooling water mass flow and temperature). This sensitivity analysis compares different modeling strategy (steady modeling, constant efficiency modeling and data driven unsteady modeling) and let to show the relevance to use such a model compared with the state of the art adopted modeling typology. The sensitivity analysis will show the importance attached mainly to the time simulation, to the electrical demand data resolution and to the level of precision of μ CHP models. About self-consumption ratios, simplified assumptions induce range differences of +10 to + 20% compared with a more detailed modeling.

References

- [1] Bianchi M., De Pascale A., Ruggero Spina P., Guidelines for residential micro-CHP systems design. *Applied Energy*, 2012, 97:673-685.
- [2] Bouvenot J.-B., Andlauer A., Stabat P., Marchio D., Flament B., Latour B., Siroux M., Gas Stirling engine μ CHP boiler experimental data driven model for building energy simulation. *Energy and Buildings*, 2014, 84:117-131.
- [3] Bouvenot J.-B., Latour B., Siroux M., Flament B., Stabat P., Marchio D., Dynamic model based on experimental investigations of a wood pellet steam engine μ CHP for building energy simulation. *Applied Thermal Eng.*, 2014, 73:1041-1054.
- [4] Bouvenot J.-B., Siroux M., Latour B., Flament B., Energetic, environmental and economic simulation platform development of μ CHP and energy storage systems coupled to buildings, *Proceeding of ECOS international conference 2015*, Pau, Juillet 2015.
- [5] Bouvenot J.-B., Siroux M., Latour B., Flament B., Dwellings electrical and DHW load profiles generators development for μ CHP systems using RES coupled to buildings applications, *Energy Procedia*, Volume(78), Pages 1919-1924, 2015.

Quantile-oriented sensitivity indices

BROWNE THOMAS^{1,2}, FORT JEAN-CLAUDE², LE GRATIET LOÏC¹

¹ EDF Lab Chatou, France

² Université Paris-Descartes, France

Goal-oriented sensitivity analysis (*GOSA*, [1])

Let f be a numerical code and Y its one-dimensional output such that $Y = f(X_1, X_2, \dots, X_d)$, where $X = (X_1, X_2, \dots, X_d)$ are independent random inputs. Regarding a certain strategy for the study, we focus on one precise property of Y 's distribution, $\theta(Y)$: it can be $\mathbb{E}[Y]$, $q^\alpha(Y)$, the α -quantile of Y , $\mathbb{P}(Y > t_s)$ with t_s a threshold. If there is a need for sensitivity analysis, *GOSA* states that it may be more relevant to restrict it to $\theta(Y)$. Therefore our wish is to quantify the inputs' influence over $\theta(Y)$. It consists in studying the variability of the conditional parameter $\theta(Y | X_i)$. We adopt the following theoretical method: for each input X_i , with $i \in \{1, \dots, d\}$, one consecutively sets $X_i = x_i$ for all the possible values of X_i and simulates $f(X_1, \dots, x_i, \dots, X_d)$ an infinite number of times. Hence one can compute $\theta(Y | X_i = x_i)$. We repeat this procedure for all the possible values x_i so that we learn $\theta(Y | X_i)$'s distribution. The latter contains the needed information about X_i 's influence over $\theta(Y)$.

Sensitivity analysis indices with respect to a contrast : the quantile case

In [2] the authors introduced the following sensitivity index, for $i \in \{1, \dots, d\}$ and $\alpha \in]0, 1[$:

$$S_{c_\alpha}^i(Y) = \min_{\theta \in \mathbb{R}} \mathbb{E}[c_\alpha(Y, \theta)] - \mathbb{E}_{X_i} \left[\min_{\theta \in \mathbb{R}} \mathbb{E}[c_\alpha(Y, \theta) | X_i] \right],$$

with: $\forall y, \theta \in \mathbb{R} \quad c_\alpha(y, \theta) = (y - \theta)(\mathbf{1}_{y \leq \theta} - \alpha)$, which evaluates the influence of X_i over $q^\alpha(Y)$. Besides, let us recall that $q^\alpha(Y) = \arg \min_{\theta \in \mathbb{R}} \mathbb{E}[c_\alpha(Y, \theta)]$ and $q^\alpha(Y | X_i = x_i) = \arg \min_{\theta \in \mathbb{R}} \mathbb{E}[c_\alpha(Y, \theta) | X_i = x_i]$, therefore:

$$S_{c_\alpha}^i(Y) = \mathbb{E}[c_\alpha(Y, q^\alpha(Y))] - \mathbb{E}[c_\alpha(Y, q^\alpha(Y | X_i))].$$

Hence one can see that $S_{c_\alpha}^i(Y)$ quantifies the modification of $q^\alpha(Y)$ when one sets X_i to a single value. Besides, we easily prove $\mathbb{E}_{X_i} \left[\min_{\theta \in \mathbb{R}} \mathbb{E}[c_\alpha(Y, \theta) | X_i] \right] \leq \min_{\theta \in \mathbb{R}} \mathbb{E}[c_\alpha(Y, \theta)]$. Then the authors normalize the index as they divide it by $\min_{\theta \in \mathbb{R}} \mathbb{E}[c_\alpha(Y, \theta)]$. This now implies: $0 \leq S_{c_\alpha}^i(Y) \leq 1$. In order to justify the meaning of the index, we prove the following property:

$$\begin{aligned} S_{c_\alpha}^i(Y) &= 0 && \text{if and only if } q^\alpha(Y | X_i) = q^\alpha(Y) \text{ a.s.} \\ S_{c_\alpha}^i(Y) &= 1 && \text{if and only if } \forall x_i \quad \text{Var}(Y | X_i = x_i) = 0. \end{aligned}$$

Estimator and property

From a n -sample $(Y^1, \dots, Y^n)$, where $n \in \mathbb{N}$, and for $j \in \{1, \dots, n\}$, $Y^j = f(X_1^j, \dots, X_d^j)$, we propose an estimator for $S_{c_\alpha}^i(Y)$. The first term can be easily estimated by a classical empirical estimation $\min_{\theta \in \mathbb{R}} \frac{1}{n} \sum_{j=1}^n c_\alpha(Y^j, \theta)$, where $\hat{q}^\alpha(Y) := \arg \min_{\theta \in \mathbb{R}} \frac{1}{n} \sum_{j=1}^n c_\alpha(Y^j, \theta)$ is the empirical quantile estimator. The second term is much more complicated to estimate as it contains a double expectation (including a conditional expectation) and requires to solve a minimization problem. We base its estimation on the following asymptotic result proved in [3]:

$$\forall x_i \text{ st } f_i(x_i) \neq 0, \arg \min_{\theta} \frac{1}{f_i(x_i)} \sum_{j=1}^n c_\alpha(Y^j, \theta) K_{h(n)}(X_i^j - x_i) \xrightarrow[n \rightarrow \infty]{\mathbb{P}} \arg \min_{\theta} \mathbb{E}[c_\alpha(Y, \theta) | X_i = x_i],$$

where K is a positive second-order kernel on a bounded compact, (h_k) the bandwidth sequence and f_i the density function of X_i , with the conditions $h(n) \xrightarrow[n \rightarrow +\infty]{} 0$ and $h(n) \times n \xrightarrow[n \rightarrow +\infty]{} +\infty$.

The problem is that we are not interested in the minimizers but in the minimal values for each possible x_i by which we condition, and need to compute their average over the different x_i . At the end, we propose the following kernel-based estimator for $S_{c_\alpha}^i(Y)$:

$$\widehat{S_{c_\alpha}^i(Y)} = \min_{\theta \in \mathbb{R}} \frac{1}{n} \sum_{j=1}^n c_\alpha(Y^j, \theta) - \frac{1}{n} \sum_{k=1}^n \min_{\theta \in \mathbb{R}} \frac{1}{k \cdot f_i(X_i^k)} \left[\sum_{j=1}^k c_\alpha(Y^j, \theta) \frac{1}{h_k} K\left(\frac{X_i^k - X_i^j}{h_k}\right) \right].$$

Under the same conditions than above we prove the consistency of the estimator:

$$\widehat{S_{c_\alpha}^i(Y)} \xrightarrow[n \rightarrow +\infty]{\mathbb{P}} S_{c_\alpha}^i(Y).$$

Applications to defect detection

We study an example in the context of defect examination: we inspect of a structure by sending a wave that reflects on the hypothetical defect. The random reflected signal Z , function of the size of defect a , random environmental properties X and a noise of observation δ , is measured so that: $(Z(a, X, \delta) > t_s)$ implies that the defect is detected. Let us focus on the random defect a_{90} , function of the inputs X , defined as $\mathbb{P}(Z(a_{90}, X, \delta) > t_s | X) = 0.90$, which is the defect that is detected with a probability of 90% under the conditions X . In this example we consider three inputs,

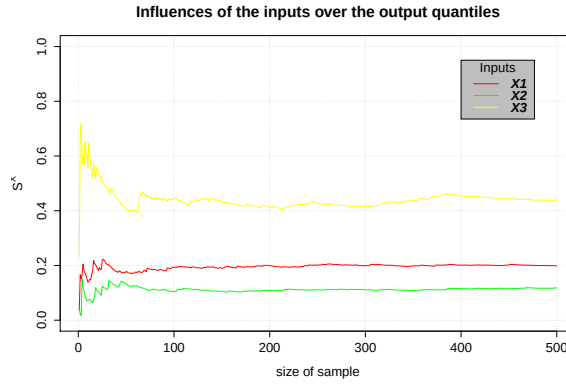


Figure 1: Index estimation for the influence of the three inputs over $q^\alpha(a_{90})$. $\hat{S}_{c_\alpha}^1(a_{90})$ is in red, $\hat{S}_{c_\alpha}^2(a_{90})$ in green and $\hat{S}_{c_\alpha}^3(a_{90})$ in yellow.

$X = (X_1, X_2, X_3)$, and the wish is to estimate $S_{c_\alpha}^1(a_{90})$, $S_{c_\alpha}^2(a_{90})$ and $S_{c_\alpha}^3(a_{90})$. The different estimators $\hat{S}_{c_\alpha}^1(a_{90})$, $\hat{S}_{c_\alpha}^2(a_{90})$ and $\hat{S}_{c_\alpha}^3(a_{90})$ are computed for a size of sample $n = 2, \dots, 150$.

Bibliographie

- [1] N. Rachdi (2011), Statistical Learning and Computer Experiments, Thèse de l'Université Paul Sabatier, Toulouse, France.
- [2] J-C. Fort, T. Klein et N. Rachdi (2013), New sensitivity analysis subordinated to a contrast, *Communication in Statistics : Theory and Methods*, In press.
- [3] J. Fan, T. Hu and Y. K. Truong (1994), Robust Non-Parametric Function Estimation, *Scandinavian Journal of Statistics*, Vol. 21, No. 4, pp. 433-446

[thomas.ga.browne@gmail.com; EDF R&D, 6 Quai Watier, 78401 Chatou, France]

Investigation of Modern Methods of Probabilistic Sensitivity Analysis of Final Repository Performance Assessment Models

D.A. Becker¹ and S.M. Spiessl¹

¹Gesellschaft fuer Anlagen- und Reaktorsicherheit (GRS) gGmbH, Germany

Abstract

For deep geological repositories for radioactive waste, numerical performance assessment is a key process in all phases from site selection to licensing for closure. The release of contaminants to the biosphere for a number of conceivable scenarios has to be assessed in advance, which can only be done by modelling all relevant effects in an integrated, coupled model. Such computation models are typically rather complex, as they combine a lot of physical and chemical effects and influences from various processes in the underground. As a result, they often show a highly non-linear behaviour.

There are many parameters influencing the calculation results that are subject to essential uncertainties. By this reason, sensitivity analysis is an important tool for investigating the model behaviour. Sensitivity analysis is not only adequate for directing research activities, but can contribute essentially to a proper model understanding and even reveal errors in the model or the data.

In the past, there was a tendency to apply well-known standard methods of probabilistic sensitivity analysis to performance assessment models uncritically without thinking about their appropriateness. Although such a procedure often leads to a correct sensitivity estimation, it cannot be excluded that, in extreme cases, it can yield wrong or misleading results and jeopardise the benefit of sensitivity analysis. Therefore, a research programme was set up some years ago in order to investigate new developments in sensitivity analysis, their applicability to performance assessment model results and the benefit such methods can provide for repository safety assessment. The final goal of the investigations was to provide some guidance to a modeller for performing an effective and meaningful sensitivity analysis. In this talk we present an overview of the total project and the main outcomes.

Three performance assessment models were defined for hypothetical repositories for different kinds of radioactive waste in different geological formations. These models show different effects that are typical for their specific type, like output results widely spread over many orders of magnitude, occurrence of a considerable number of zero - runs, a two - split output distribution or an extremely non-linear, nearly non-continuous behaviour. For each model a set of uncertain input parameters was defined. An appropriate pdf was assigned to each parameter.

The models were calculated a high number of times using parameter samples of sizes between 1000 and 32000 that were drawn applying different sampling algorithms like Random sampling, Latin Hypercube sampling, Quasi-Random-LpTau sampling, FAST and EFAST sampling, and Random Balance Design (RBD) sampling. Different methods of sensitivity analysis were applied, including Standardised Regression and Rank Regression Coefficients (SRC/SRRC), (Extended) Fourier Amplitude Sensitivity Test (FAST/EFAST), Effective Algorithm for Computing Global Sensitivity Indices (EASI), the State-Dependent Parameter (SDP) method as well as the Smirnov test. Some experiments were also done with correlated input parameters and transformation of model output. Moreover, graphical methods

of sensitivity analysis, mainly the Contribution to Sample Mean (CSM) plot, were applied.

Sensitivity measures were calculated with each method for a number of points in time, so that the results could be plotted as time curves. The investigation of the results was oriented at the following questions:

- How robust are the results? Do the curves considerably change if a different sample of same size is used? How many runs are necessary to achieve stable curves?
- Do the different methods calculating variance-based sensitivity indices of first order produce similar results?
- Do the different sensitivity measures and graphical methods qualitatively agree about the main sensitivities?
- Are the sensitivity analysis results plausible and understandable?
- Are all sensitivities detected by the different methods?
- Which sampling algorithm seems best?
- Can the significance of sensitivity analysis be improved by transforming the model output to a more appropriate scale?
- How numerically effective are the different methods of sensitivity analysis?

Acknowledgement: This work was funded by the German Federal Ministry for Economic Affairs and Energy (BMWi) under grant No. 02E10941.

Optimising Composite Indicators with Sensitivity Analysis

William Becker¹, Michaela Saisana¹, Paolo Paruolo¹, Ine Vandecasteele² and Andrea Saltelli^{3,4}

¹European Commission, Joint Research Centre, Deputy Director-General, Econometrics and Applied Statistics Unit, Via E Fermi 2749, Ispra (VA), Italy

²European Commission, Joint Research Centre, Institute for Environment and Sustainability, Sustainability Assessment Unit, Via E Fermi 2749, Ispra (VA), Italy

³Centre for the Study of the Sciences and the Humanities (SVT), University of Bergen (UIB), Spain

⁴Institut de Ciència i Tecnologia Ambientals (ICTA), Universitat Autònoma de Barcelona, Spain

Email: william.becker@jrc.ec.europa.eu

Multi-dimensional measures (often termed composite indicators) are popular tools in the public discourse for assessing the performance of countries/entities on human development, perceived corruption, innovation, competitiveness, or other complex phenomena that are not directly measurable and not precisely defined. These measures combine a set of relevant variables using an aggregation formula, which is often a weighted arithmetic average. The values of the weights are usually meant to reflect the variables' importance in the index, which is based on the subjective beliefs of the developer. In practice, however, correlations between variables mean that the weights assigned to each variable do not actually reflect the true importance in terms of their contribution to the composite indicator. To elaborate, let $\{x_i\}_{i=1}^d$ be the set of d input variables to the composite indicator, and y be the output (i.e. the composite indicator value). Weights w_i are assigned such that:

$$y = w_1x_1 + w_2x_2 + \dots + w_dx_d,$$

where $\sum_{i=1}^d w_i = 1$. Given a sample of N points, consider an importance measure I which measures the influence of each x_i on y , which is also normalised to sum to 1. The key point is that $I_i \neq w_i$, nor is I necessarily linearly related to w , although this fact is sometimes overlooked by developers. Note that the importance of a composite indicator's inputs on its outputs is dependent on the sample. Two questions immediately arise: first, given a set of weights and a sample, what is the influence of each variable on the output? Second, how can weights be assigned to reflect the desired importance?

This work views the problem from a sensitivity analysis perspective, using tools from the literature on global sensitivity analysis with correlated inputs to understand the importance of each variable's contribution to the index. In particular, the first order sensitivity index, $S_i = \text{var}[E(Y|X_i)]/\text{var}(Y)$ is used, which is referred to here as the *Pearson correlation ratio* (an equivalent term that was used by Karl Pearson first in 1905). This term is used to emphasise that the measure is being used first and foremost to measure correlation, and not sensitivity in terms of a variance decomposition.

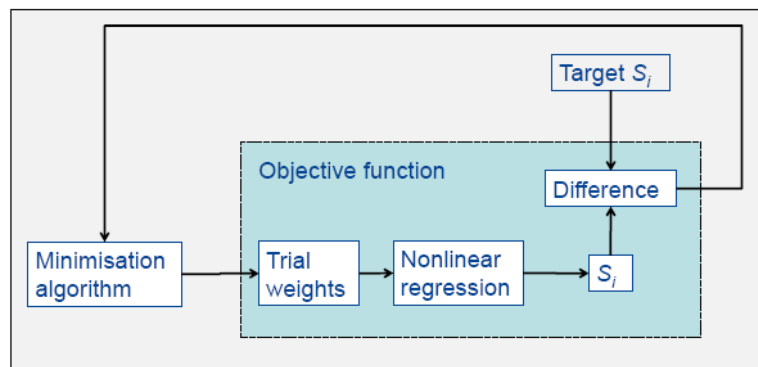
In order to estimate the correlation ratio, this work follows previous work of [3] and [1] by using nonlinear regression to estimate the main effect $E(Y|X_i)$. However, additional to the use of local polynomial regression, two other approaches are considered. The first is the use of *penalised splines*, which can be fit with a particularly low computational cost (an advantage which is exploited in the optimisation step below). The second is the use of Bayesian *Gaussian processes*, which have the advantage of providing confidence intervals on the Pearson correlation ratio.

As a further step, to better understand the influence of variables on the composite indicator, an approach based on the correlated sensitivity analysis work of [4] is applied, which uses an additional regression to decompose the influence of each variable into influence caused by correlation, and influence caused by the composite indicator structure (aggregation and weights).

To now address the second question, the issue of optimisation of weights is considered. Although this problem has been tackled in [3] using linear regression, the proposal here is to extend it to nonlinear regression, to account for nonlinear main effects. Letting $\tilde{S}_i$ be the desired correlation ratio of variable x_i , the set of weights $\mathbf{w}_{\text{opt}}$ that minimises the difference between $\tilde{S}_i$ and $S_i(\mathbf{w})$ is found by,

$$\mathbf{w}_{\text{opt}} = \text{argmin} \sum_{i=1}^d (\tilde{S}_i - S_i(\mathbf{w})),$$

where $\mathbf{w} = \{w_i\}_{i=1}^d$. This minimisation problem is performed by the Nelder-Mead simplex search method [2]. See the figure below for an overview of the optimisation process.



The methodologies proposed here are applied to several test cases, namely the Resource Governance Index, the Good Country index, and a hydrological example—the Water Retention index, which demonstrates how weight optimisation can be performed on composite indicators with thousands, or possibly millions, of data points. The case studies provide insight in terms of the ideal weightings for each composite indicator, as well as illustrating the potential and limitations of the proposed approaches.

[1] Da Veiga, S., Wahl, F., & Gamboa, F. (2009). Local polynomial estimation for sensitivity analysis on models with correlated inputs. *Technometrics*, 51(4), 452-463.

[2] Lagarias, J. C., Reeds, J. A., Wright, M. H., & Wright, P. E. (1998). Convergence properties of the Nelder–Mead simplex method in low dimensions. *SIAM Journal on optimization*, 9(1), 112-147.

[3] Paruolo, P., Saisana, M., & Saltelli, A. (2013). Ratings and rankings: voodoo or science?. *Journal of the Royal Statistical Society: Series A (Statistics in Society)*, 176(3), 609-634.

[4] Xu, C., & Gertner, G. Z. (2008). Uncertainty and sensitivity analysis for models with correlated parameters. *Reliability Engineering & System Safety*, 93(10), 1563-1573.

A Copula-based Approach to Sensitivity to Correlations in Structural Reliability Problems

NAZIH BENOUMECHIARA
LSTA-UPMC & EDF Lab Chatou, France

ROMAN SUEUR & NICOLAS BOUSQUET & BERTRAND IOOSS
EDF Lab Chatou, France

GÉRARD BIAU & BERTRAND MICHEL & PHILIPPE SAINT-PIERRE
LSTA-UPMC & Institut de Mathématique de Toulouse, France

To ensure the high reliability level of industrial structures, EDF conducts probabilistic studies [1]. They are based on a computational model, which aims at describing at best the physical behaviour of a structure under loading. A statistical model is built to describe the uncertainties of the parameters involved in the computational model. Unfortunately, little information is usually available on the stochastic dependence of variables. The statistical model is therefore partial and can be reduced to its margins only. Consequently, reliability studies in industrial practice are frequently carried out assuming independence of variables. A question that arises is how can we enhance the robustness of the studies without knowledge of the correlations? To answer this, our work aims to quantify the impact of potential dependencies on the structure reliability.

The methodology considers the input random vector $\mathbf{X} = (X_1, \dots, X_d) \in \mathcal{S}_{\mathbf{X}}$ and the output random variable $Y = g(\mathbf{X}) \in \mathcal{S}_Y$ of the model g . The quantity of interest of the output variable Y , used to quantify the risk faced by the structure, is denoted by $\mathcal{C}(Y)$. We use the notion of copulas to describe the dependence structure of $\mathbf{X}$, independently of its marginals. The joint Cumulative Distribution Function (CDF) of $\mathbf{X}$ is thus given as

$$F_{\mathbf{X}}(x_1, \dots, x_d) = C_{\boldsymbol{\rho}}(F_{X_1}(x_1), \dots, F_{X_d}(x_d)),$$

where $C_{\boldsymbol{\rho}} : [0, 1]^d \rightarrow [0, 1]$ is a copula with parameter $\boldsymbol{\rho} \in S_{\boldsymbol{\rho}}$ and F_{X_i} is the marginal's CDF of X_i . We also introduce the notation $\mathbf{X}^{\boldsymbol{\rho}}$ to describe a random vector $\mathbf{X}$ associated with a copula $C_{\boldsymbol{\rho}}$, and the related output variable $Y^{\boldsymbol{\rho}} = g(\mathbf{X}^{\boldsymbol{\rho}})$.

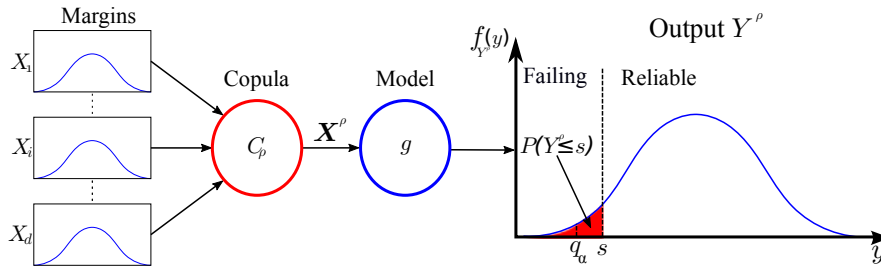


Figure 1: Uncertainty propagation of $\mathbf{X}$ with a copula $C_{\boldsymbol{\rho}}$ through the model g .

Some related studies focused on measuring the impact of a perturbation on a marginal X_i [2] or an incomplete joint density of $\mathbf{X}$ [3] on the model output Y . In this work, we propose a sensitivity index which quantifies, for a chosen copula, the change on the quantity of interest between the worst case scenario and the independence case. Such an index is described by

$$\mathcal{I} = \frac{\mathcal{C}(Y^{\boldsymbol{\rho}^*})}{\mathcal{C}(Y)},$$

where $\mathcal{C}(Y)$ is the quantity of interest of Y at independence and ρ^* is the dependence configuration obtained by maximising the risk R , such as $\rho^* = \operatorname{argmax}_{\rho \in \mathcal{S}_\rho} R(\rho)$. The index $\mathcal{I}$ would describes the general impact of the dependence structure on $\mathcal{C}$. Moreover, the index $\mathcal{I}_{ij}$ quantifies the impact of a one pair of variables dependence X_i-X_j on $\mathcal{C}$, while the other variables are independent. As for the general index $\mathcal{I}$, it is defined as

$$\mathcal{I}_{ij} = \frac{\mathcal{C}(Y^{\rho_{ij}^*})}{\mathcal{C}(Y)},$$

where $\rho_{ij}^* = \operatorname{argmax}_{\rho \in \mathcal{S}_\rho} R(\rho_{ij})$ is the dependence parameter of the pair of variables X_i-X_j leading to the worst case scenario. There is, for the moment, no direct relation between $\mathcal{I}_{ij}$ and $\mathcal{I}$.

The estimation of such indices is almost entirely controlled by the estimation of the worst case dependence structure ρ^* . This problem of *extremum-estimation* is consistent using a Monte-Carlo sampling. The figure 2 shows the related estimated one pair indices, for a Gaussian copula, applied to the Flood example [4] using the failure probability of Y as the quantity of interest. The closer an index $\mathcal{I}_{ij}$ is to 1 and the less impactful the dependence of the pair X_i-X_j is. And vice versa when the value is high.

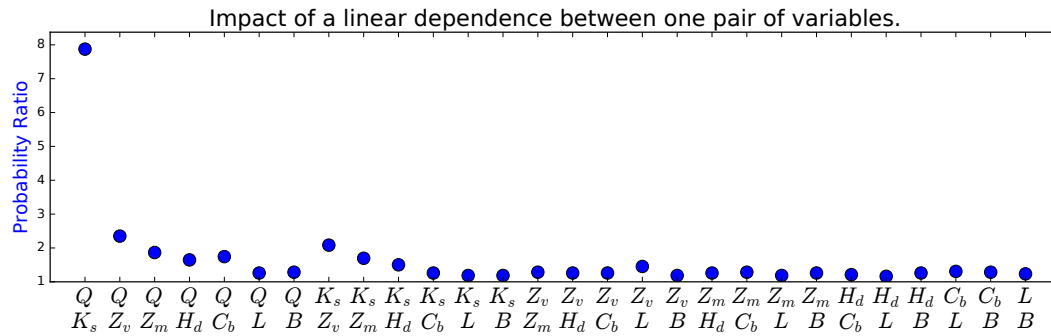


Figure 2: Monte-Carlo estimation of the indices $\mathcal{I}_{ij}$ for each pair X_i-X_j of the flood example.

Unfortunately, the Monte-Carlo sampling is costly and can hardly be performed for computationally expensive models. Thus, other estimation methods, such as Random Forests [5], could be considered to reduce the number of model evaluations. Moreover, such indices can be too pessimistic, because the worst case scenario can be very unlikely. Therefore, another perspective would be to consider every penalised dependence structures instead of the worst case configuration only.

References:

- [1] de Rocquigny, E., Devictor, N., & Tarantola, S. (Eds.). (2008). *Uncertainty in industrial practice: a guide to quantitative uncertainty management*. John Wiley & Sons.
- [2] Lemaître, P., Sergienko, E., Arnaud, A., Bousquet, N., Gamboa, F., Iooss & B. (2015). Density modification-based reliability sensitivity analysis. *Journal of Statistical Computation and Simulation*, 85(6), 1200-1223.
- [3] Agrawal, S., Ding, Y., Saberi, A., & Ye, Y. (2012). Price of correlations in stochastic optimization. *Operations Research*, 60(1), 150-162.
- [4] B. Iooss & P. Lemaître. A review on global sensitivity analysis methods. In: *Uncertainty management in Simulation-Optimization of Complex Systems: Algorithms and Applications*, C. Meloni and G. Dellino (Eds.), Springer, 2015.
- [5] Breiman, L. (2001). Random forests. *Machine learning*, 45(1), 5-32.

Sensitivity analysis under different distributions using the same simulation

R. Bolado Lavin¹ and S. Tarantola¹

¹Directorate for Energy, Transport and Climate, European Commission, Joint Research Centre

Abstract

Modelling is used in many areas of science and technology to deal with safety and security problems. When dealing with this type of problems, risk / safety / security criteria are established and the technical body of the organization developing the study has to prove compliance of the system under study with the criteria (think, for example, of nuclear reactor safety, high level nuclear radioactive waste repositories safety assessments or security of gas supply). Criteria are typically established on some output variable, such as for example the dose to the population or the quantity of unserved gas. Most frequent criteria are based either on the expected value of the output variable or on some probability of exceedance (the expected dose should not exceed a given reference value, or the probability of not satisfying the demand of gas protected customers at least one day of the year should be less than 5%).

In many of these problems, the uncertainty in the input parameters is characterized by means of probability density functions (pdf) via expert judgement, given the intrinsic difficulty, even impossibility of taking actual measurements. When this happens, sometimes Sensitivity Analysis (SA) practitioners and analysts are confronted with the problem of having different, conflicting pdfs for characterizing the uncertainty in a given input parameter. Sometimes experts acknowledge large uncertainty in the scale and shape of the input parameters; sometimes they provide conditional probabilities (dependent on other uncontrolled parameters); sometimes they even do not agree. Even if pdfs are obtained via actual measurements, pdfs may evolve over time after the acquisition of new information (Bayesian update of information). Under these circumstances, SA practitioners are asked how the output variables distributions could change given changes in the input parameters pdfs, and in particular, if alternative input parameters pdfs could deliver significant changes in terms of criteria violation.

During the last two decades we have seen a huge development in the area of SA (variance based techniques, Monte Carlo filtering, graphical techniques, etc.), but this problem has been systematically ignored in the literature. The trivial solution to this problem is to execute Monte Carlo with the default input multivariate distribution ($f_1(\mathbf{x})$) and to repeat it again, using a new sample, using the alternative ($f_2(\mathbf{x})$) pdf, comparing afterwards the statistics obtained and apply the corresponding test to determine if differences are statistically significant (provided that the adequate test exists). Certainly this procedure is far from optimal. McKay and Beckman [1] are among the few authors that have addressed formally this problem. These authors propose two methods to estimate the effect of changing the input distributions: the rejection method and the weighting method. Hesterberg [2] addresses also the problem from the point of view of importance sampling. The approach adopted in this paper is not the importance sampling approach, but the SA view. The estimator proposed for the distribution of the output variable is the one that assigns a new weight

$$P(Y(\mathbf{x}_i)) = \frac{\frac{f_2(\mathbf{x}_i)}{f_1(\mathbf{x}_i)}}{\sum_{j=1}^n \frac{f_2(\mathbf{x}_j)}{f_1(\mathbf{x}_j)}}$$

to each sampled output value Y . In this way one can estimate the pdf of the output for different pdfs of the input and, in particular, variance-based sensitivity indices for different input pdfs, using the same simulations.

We have applied this method to the assessment of the safety of a passive system in a nuclear power plant. This system is the BOPHR/RP2 "Base Operation Passive Heat Removal strategy applied to Residual Passive heat Removal system on the Primary circuit. This system is considered passive because after the start of the nuclear accident it relies only of physical principles to accomplish its mission (adequately cool down the reactor core and keep moderate the pressure in the primary circuit). In particular it should be able to work with total lack of electricity supply, relying only on natural circulation. In this problem 14 input parameters affected by uncertainty were considered. The system was considered to succeed if it were able to keep pressure in the primary circuit below 4 MPascal.

Figures 1 and 2 show the results obtained for the Pressure when the pdfs of two different parameters are changed. Input parameter 1 is very important, while input parameter 2 is irrelevant to the output variable considered. Under the default distribution (purple line), the estimated probability of failure of the system (pressure exceeding 4MPascal) is approx. 0.04. Under the different alternative pdfs for parameter 1, the probability of failure may change between approx. 0.01 and 0.11. On the other hand, the probability is completely insensitive to changes in the pdf of input parameter 2. In the paper we show how to derive the proposed approach, its properties and statistical tests applicable to determine if the differences are statistically significant.

References

1. Beckman R.J., McKay M.D. Monte Carlo Estimation Under Different Distributions Using the Same Simulation. *Technometrics* Vol 29, No 2, pages 153-160. 1987.
2. Hesterberg T. Weighted Average Importance Sampling and Defensive Mixture distributions. *Technometrics* Vol 37, No 2, pages 185-195. 1995.

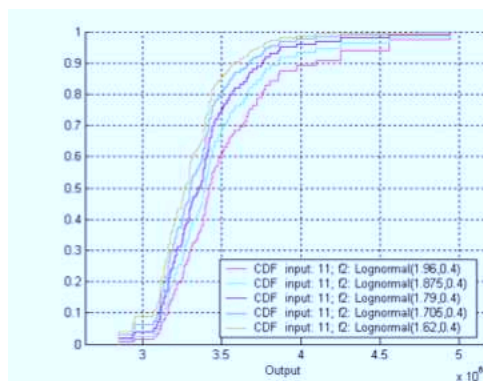


Fig. 1. Different empirical distributions of the output variable for different distributions of input parameter 1.

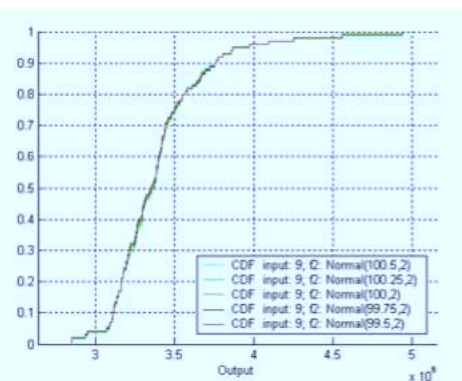


Fig. 1. Different empirical distributions of the output variable for different distributions of input parameter 2.

Predicted Sensitivity for establishing well-posedness conditions in stochastic inversion problems

MÉLANIE BLAZÈRE & NICOLAS BOUSQUET

Institut de Mathématique de Toulouse & EDF Lab Chatou, France

We consider a stochastic inversion problem defined by the knowledge of observations $\mathbf{y}_n = (y_i^*)_{i \in \{1, \dots, n\}}$ living in a q -dimensional space, which are assumed to be realizations of a random variable Y^* such that

$$\begin{aligned} Y^* &= Y + \varepsilon, \\ Y &= g(X) \end{aligned}$$

where X is a d -dimensional random Gaussian variable $X \sim \mathcal{N}(\mu, \Sigma)$ with unknown $\theta = (\mu, \Sigma)$, ε is a (experimental or/and process) noise with known distribution f_ε , and g is some deterministic function from $\mathbb{R}^p$ to $\mathbb{R}^q$ (possibly a black-box computer model). This inversion problem (ie., estimating θ) can be solved in frequentist [3,1] or Bayesian [4] frameworks (possibly by linearizing g [1]), using missing data algorithms. In both frameworks, inferring on θ requires that several conditions of well-posedness and identifiability are gathered.

The first one is Hadamard's well-posedness condition, which states that the solution $\hat{\theta}$ of the inversion/calibration problem should exist, be unique and be continuously dependent on observations according to a reasonable topology. In the case where g is linear or can be linearized, namely if there exists a linear operator H such that $Y^* = HX + \varepsilon$, this condition is traduced by a low value of the *condition number* of H [2]. The second condition is the identifiability of the input model $X \sim \mathcal{N}(\mu, \Sigma)$. In similar cases of linearity or linearization, this condition states that H must be injective ($\text{rank}(H) = d$) and $d \leq nq$.

However, additionally to Hadamard's condition, and independently of the availability of experimental data $\mathbf{y}^*$, a second condition of well-posedness is, to our knowledge, never evoked in practice, while it seems to be of primary importance in the specific framework of stochastic inversion. This condition arises from (let us say) *predictive sensitivity analysis*. Imagine that the problem is solved and θ is known. Any sensitivity study, for instance based on celebrated Sobol' indices [5], should highlight that the main source of uncertainty, explaining the variations of Y^* , is X and not ε . In practice, this kind of diagnostic is established a posteriori, as a check for an estimated solution $\hat{\theta}$ (or a posterior distribution $\pi(\theta|\mathbf{y}_n)$ in a Bayesian context). However, this property is more than desirable and should be converted into a modelling constraint for the estimation of θ . In a Bayesian inversion context, such a constraint would apply on the prior elicitation of (the parameters of) θ , and could help to define better reference measures when no other prior information is available on θ nor X . Such a study requires a formal definition of what "the main source of uncertainty" means.

Several answers to the problem of well-defining a stochastic inversion problem, by formalizing a new condition on θ with respect to the features of g and f_ε , are proposed in this talk. They are successively based on Sobol' indices, entropic indices then comparisons of Fisher information. The case when g is linear or linearizable is considered in this study, since it appears as a minimal framework to establish such a rule (or possibly several rules) of well-posedness. Strongly nonlinear cases often by definition present more contrasted behavior between observations and noise, and it is likely that the ratio between signal and noise be more in favor of the signal.

In this regard, one-dimensional linear then linearizable models are studied. Then a general rule of well-posedness is presented and results are derived for multivariate linearizable models. A theoretical link is done with Sobol'indices. Extensions to nonlinear cases are discussed, as well as stochastic metamodels as Gaussian kriging used in general stochastic inversion [4]. Furthermore, the control of bias arising in linearizable contexts is evoked, as a new source of uncertainty that should be monitored in the same way as the noise ε . Finally, the approach is tested over toy examples and a simplified hydraulical example, and compared with usual stochastic inversion methodologies that do not consider this well-posedness condition a priori.

References:

- [1] P. Barbillon, G. Celeux, A. Grimaud, Y. Lefebvre, and E. Rocquigny (de). "Non linear methods for inverse statistical problems". *Computational Statistics and Data Analysis* 55, 132-142, 2011.
- [2] D.A. Belsley, E. Kuh, and R.E. Welsch. "The Condition Number". In: *Regression Diagnostics: Identifying Influential Data and Sources of Collinearity*. New York: John Wiley & Sons, 1980.
- [3] G. Celeux, A. Grimaud, Y. Lefebvre, and E. Rocquigny (de). "Identifying intrinsic variability in multivariate systems through linearised inverse methods". *Inverse Problems in Engineering* 18, 401-415, 2010.
- [4] S. Fu, G. Celeux, N. Bousquet, and M. Couplet. "Bayesian inference for inverse problems occurring in uncertainty analysis". *International Journal for Uncertainty Quantification* 5, 73-98, 2015.
- [5] B. Iooss and P. Lemaître. "A review on global sensitivity analysis methods". In: *Uncertainty management in Simulation-Optimization of Complex Systems: Algorithms and Applications*, C. Meloni and G. Dellino (eds.), Springer, 2015

Sensitivity analysis and the calibration problem of a biodynamic model for Indian Ocean tuna growth

NICOLAS BOUSQUET¹ & EMMANUEL CHASSOT² & SÉBASTIEN DA VEIGA³ & THIERRY KLEIN¹ & BERTRAND IOOSS¹ & AGNÈS LAGNOUX¹

Institut de Mathématique de Toulouse¹ & Institut de Recherche pour le Développement² & SAFRAN Tech³, France

The growth of Indian Ocean Yellowfin (*Thunnus albacares*) tunas is characterized by several shapes (Figure 1) which can be explained by the juxtaposition of several *environmental forcing* parameters $X \in \Omega \subset \mathbb{R}^d$ (for instance the water temperature or the food density) and *metabolic* (intrinsic) parameters $\theta \in \Theta \subset \mathbb{R}^q$. Several regime shifts can appear, that are connected to specific development stages: typically, larvae grow fast to juveniles, then some bioenergetic losses may occur, indicating spawning or senescence. Nonetheless, older and healthy fish are experimented predators, which may be traduced by a sensible increase of growth acceleration towards an asymptotic limit, in favorable conditions. The knowledge of the values range of the most influential parameters driving each kind of shape and the age-length key also produced can play an important role in the determination of the size structure of fishing gears, in a perspective of elaborating sustainable exploitation patterns.

Calibrating and classifying the range of values for (X, θ) , in function of shapes, can be conducted using a biodynamical, functional computer model $\{L(t), W(t)\}_{t=0, \dots, T} = g_\theta(X)$ based on the Dynamic Energy Budget (DEB) theory [1]. It simulates simultaneously the curves of fork length $L(t)$ and the weight $W(t)$ indexed by age $0 \leq t \leq T$, and several sources of information. The latter are described as pointwise noisy observations $\mathbf{D}^*$ of fork lengths and weights arising from laboratory experiments, commercial catches and capture-recapture campaigns (Figure 2). More formally, the aims are:

1. to start from two prior distributions $f_X(x)$ and $\pi(\theta)$, typically based on previous works on close species, and to conduct a first sensitivity analysis (SA) to highlight the most influent inputs and decrease the dimension ($d+q \geq 20$); classical SA tools (differential SA, multivariate Sobol' indices [2], etc., see [6] for a review) and others (e.g., elasticity indices) are used to do so, since no dependence between the inputs is assumed a priori;
2. to update $f_X(x)$ and $\pi(\theta)$ conditionally to the likelihood $\ell(\mathbf{D}^*|X, \theta)$ by computing the posterior distribution with density

$$h(x, \theta|\mathbf{D}^*) \propto \ell(\mathbf{D}^*|x, \theta)f(x)\pi(\theta); \quad (1)$$

3. to conduct a more detailed sensitivity study, taking into account the classification of shapes arising from the knowledge of (1) and the correlation between the inputs. More elaborated SA approach are required, as in [3].

The Bayesian computation of the posterior faces difficulties linked to the nature of observations: capture-recapture data have correlated noises, while one of the dataset $\mathbf{D}_3 \subset \mathbf{D}$ provides a huge number of correlated fork lengths and weights, without age index. Therefore the full likelihood of observations is not tractable.

Hence, in a first step, “likelihood-free” methods as Approximate Bayesian Computation (ABC, [4]) appear to be useful to update the original prior $f_X(x)\pi(\theta)$ in a more informational prior $h_\epsilon(x, \theta) \propto \mathcal{M}\{\mathbf{D}_3^*, \mathbf{D}_3(x, \theta)\}f_X(x)\pi(\theta)$, where $\mathcal{M}\{.,.\}$ is a set of similarity measures between true ($\mathbf{D}_3^*$) and simulated ($\mathbf{D}_3(x, \theta)$) observations. It is indexed by a parameter vector ϵ determining a limit value for the similarity. This ABC step is followed by a kernel reconstruction of the new prior $\tilde{h}_\epsilon(x, \theta)$ based on copulas (R-Vines) [5]. This formal prior distribution, indeed, allows to

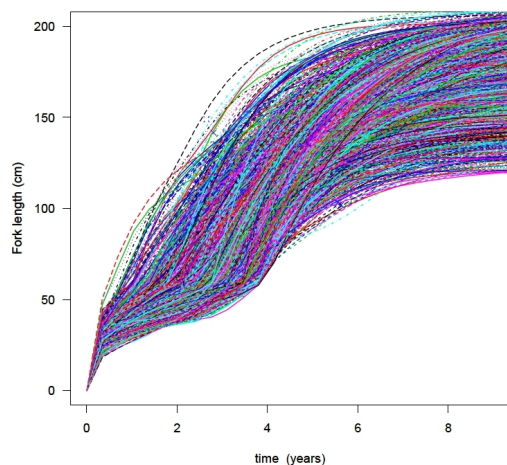


Figure 1: Typical shapes of Indian Ocean Yellowfin growth (fork length versus age).

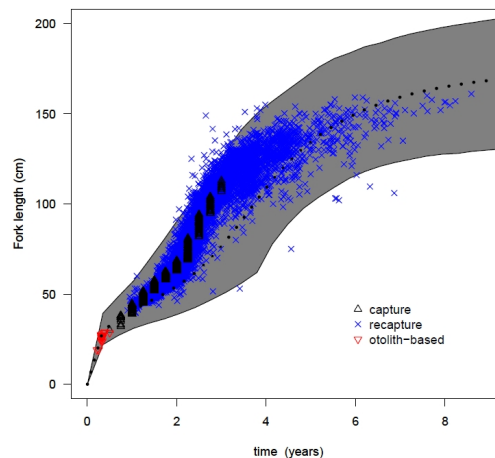


Figure 2: Posterior credibility (90%) area for growth curves arising from the calibrated model (fork length versus age).

develop an adaptive algorithm of posterior calibration based on classic Markov Chain Monte Carlo (MCMC) methods.

The results highlight the strong mixed several influence of environmental and intrinsic parameters, and the calibration methodology developed in this work allows to discriminate the influential factors that explain the variety of shapes. This can apply to a large number of studies involving noisy ground experiments and difficulties to connect correlated observations to functional computer models through statistical likelihoods.

References:

- [1] E. Dortel, L. Pecquerie, E. Chassot. "A Dynamic Energy Budget modeling approach to investigate the eco-physiological factors behind the two-stanza growth of yellowfin tuna". *Submitted*, 2016.
- [2] F. Gamboa, A. Janon, Th. Klein, A. Lagnoux. "Sensitivity analysis for multidimensional and functional outputs". *Electronic Journal of Statistics*, 8, 575-603, 2014.
- [3] S. Da Veiga. "Global sensitivity analysis with dependence measures". *Journal of Statistical and Computational Simulation*, 85, 1283-1305, 2015.
- [4] J.-M. Marin, P. Pudlo, C.P. Robert, R.J. Ryder. "Approximate Bayesian computational methods". *Statistics and Computing*, 22, 1167-1180, 2012.
- [5] T. Bedford, R. M. Cooke. (2001). "Probability density decomposition for conditionally dependent random variables modeled by vines". *Annals of Mathematics and Artificial Intelligence* 32, 245-268, 2001.
- [6] B. Iooss and P. Lemaître. "A review on global sensitivity analysis methods". In: *Uncertainty management in Simulation-Optimization of Complex Systems: Algorithms and Applications*, C. Meloni and G. Dellino (eds.), Springer, 2015

Link between the sensitivity indices at different scales

Bruchou, C. ^{*} and Saint-Geours, N. [†]

^{*}Biosp INRA

84914 Avignon, France

[†] itk

CEEI CAP ALPHA - Avenue de l'Europe, 34830 Clapiers, France

1 Objectives

It is well known that importance of the impact of input factors of a numerical model may depends on the range of variation of factors. That importance can also vary by regions in the domain of definition. Local and Global sensitivity analyses evaluate factor impacts at a point or along the whole range of the domain respectively (Saltelli et al, 2004). Sobol and Kucherenko (2009) defined the derivative based indices and showed their link with global ones. The first objective of that work in progress is to show the link between variance based sensitivity indices at different scales. The scales are defined by cutting the domain of definition. The second objective, that is a practical one, aims to estimate the terms of the relationship linking the indices. A method of estimation is proposed and evaluated on an analytical deterministic model.

2 Notations and definitions

- f a real square-integrable function defined on hypercube $\Omega = [0, 1]^K$, $\mathcal{F} = \{1, \dots, K\}$:
 - $\mathbf{X}^{(\mathcal{F})} = (X^{(1)}, \dots, X^{(K)})$ the vector of the variates or factors of f , $\mathbf{x}$ a value of $\mathbf{X}^{(\mathcal{F})}$,
 - $\mathcal{F} = \mathcal{I} \cup \{\sim \mathcal{I}\}$, with $\mathcal{I} = \{1, \dots, M\}$ and $\sim \mathcal{I}$ complement of $\mathcal{I}$ in $\mathcal{F}$,
 - $\mathbf{X}^{(\mathcal{F})} = (\mathbf{X}^{(\mathcal{I})}, \mathbf{X}^{(\sim \mathcal{I})})$, $\mathbf{x} = (\mathbf{x}_i, \mathbf{x}_j)$ with $\mathbf{x}_i \in \Omega_{\mathcal{I}} = [0, 1]^M$ and $\mathbf{x}_j \in \Omega_{\sim \mathcal{I}} = [0, 1]^{K-M}$,
- a tiling of $\Omega_{\mathcal{F}}$:
 - $\forall k \in \mathcal{F}$, $[0, 1] = \bigcup_{q \in \mathcal{Q}} I_q^{(k)}$, $\mathcal{Q} = \{1, \dots, Q\}$, with $|I_q^{(k)}| = \delta_q^{(k)}$, a partition of $[0, 1]$,
 - $\omega_i^{(\mathcal{I})} = \prod_{l=1}^M I_{i_l}^{(l)}$ cuboid in $\Omega_{\mathcal{I}}$, having volume $\delta_i = \prod_{l=1}^M \delta_{i_l}^{(l)}$,
 - $\omega_{i,j} = \omega_i^{(\mathcal{I})} \times \omega_j^{(\sim \mathcal{I})} \in \Omega$, $\mathbf{i} \in \mathcal{Q}^M$, $\mathbf{j} \in \mathcal{Q}^{K-M}$ cuboid in Ω , $\Omega = \bigcup_{\mathbf{i} \in \mathcal{Q}^K} \omega_i^{(\mathcal{F})}$.
- statistics:
 - $\mu_{i,j} = \frac{1}{\delta_i \delta_j} \int_{\omega_i^{(\mathcal{I})}} \int_{\omega_j^{(\sim \mathcal{I})}} f(\mathbf{x}_i, \mathbf{x}_j) d\mathbf{x}_j d\mathbf{x}_i$, mean of f in $\omega_{i,j}$, μ_{Ω} mean of f in Ω ,
 - $\sigma_{i,j}^2 = \frac{1}{\delta_i \delta_j} \int_{\omega_i^{(\mathcal{I})}} \int_{\omega_j^{(\sim \mathcal{I})}} (f(\mathbf{x}_i, \mathbf{x}_j) - \mu_{i,j})^2 d\mathbf{x}_j d\mathbf{x}_i$ variance of f in $\omega_{i,j}$, σ_{Ω}^2 variance of f in Ω ,
 - part of the variance of the regression of $f(\mathbf{X})$ to $\mathbf{X}^{(\mathcal{I})}$ in Ω (resp. $\omega_{i,j}$) to σ_{Ω}^2 (resp. $\sigma_{i,j}^2$) :

$$\mathbf{P}_{\Omega}^{(\mathcal{I})} = \frac{\mathbb{V}_{\Omega_{\mathcal{I}}}(\mathbb{E}_{\Omega_{\sim \mathcal{I}}}(f(\mathbf{X})|\mathbf{X}^{(\mathcal{I})}))}{\sigma_{\Omega}^2} \text{ (resp. } \mathbf{P}_{\omega_{i,j}}^{(\mathcal{I})} = \mathbb{V}_{\omega_i^{(\mathcal{I})}}(\mathbb{E}_{\omega_j^{(\sim \mathcal{I})}}(f(\mathbf{X})|\mathbf{X}^{(\mathcal{I})})) \text{)},$$
 - $\bar{f}$ the function defined on grid $\mathcal{Q}^K$: $\bar{f}(\mathbf{u}) = \frac{1}{\delta_{\mathbf{u}}} \int_{\omega_{\mathbf{u}}} f(\mathbf{x}) d\mathbf{x}$, $\mathbf{u} \in \mathcal{Q}^K$,
 - $\tilde{\mathbf{P}}_{\mathcal{Q}^K}^{(\mathcal{I})} = \frac{\mathbb{V}_{\mathcal{Q}^K}(\mathbb{E}_{\mathcal{Q}^{K-M}}(\bar{f}(\mathbf{U})|\mathbf{U}^{(\mathcal{I})}))}{\sigma_B^2}$, where $\sigma_B^2 = \sum_{\mathbf{u} \in \mathcal{Q}^K} \delta_{\mathbf{u}} (\bar{f}(\mathbf{u}) - \mu_{\Omega})^2$.

3 Relationship between the sensitivity indices at two scales

Let $m_{i,j'}^{(\mathcal{I})}(\mathbf{x}_i) = \frac{1}{\delta_{j'}} \int_{\omega_{j'}^{(\mathcal{I})}} f(\mathbf{x}_i, \mathbf{x}_j) d\mathbf{x}_j$ (The exponent $(\mathcal{I})$ highlights the space of definition $\Omega_{\mathcal{I}}$ of m). The covariance of $m_{i,j'}^{(\mathcal{I})}$ and $m_{i,j''}^{(\mathcal{I})}$ is : $\mathbb{C}(m_{i,j'}^{(\mathcal{I})}, m_{i,j''}^{(\mathcal{I})}) = \frac{1}{\delta_i} \int_{\omega_i^{(\mathcal{I})}} \left(m_{i,j'}^{(\mathcal{I})}(\mathbf{x}_i) - \mu_{i,j'} \right) \left(m_{i,j''}^{(\mathcal{I})}(\mathbf{x}_i) - \mu_{i,j''} \right) d\mathbf{x}_i$.

The index $\mathbf{P}_{\Omega}^{(\mathcal{I})}$ verifies:

$$\mathbf{P}_{\Omega}^{(\mathcal{I})} = \underbrace{\sum_{\substack{i \in \mathcal{Q}^M \\ j \in \mathcal{Q}^{K-M}}} \delta_i (\delta_j)^2 \frac{\sigma_{i,j}^2}{\sigma_{\Omega}^2} \mathbf{P}_{\omega_{i,j}}^{(\mathcal{I})}}_A + \underbrace{\frac{\sigma_B^2}{\sigma_{\Omega}^2} \tilde{\mathbf{P}}_{\mathcal{Q}^K}^{(\mathcal{I})}}_B + \underbrace{\sum_{\substack{i \in \mathcal{Q}^M \\ j', j'' \in \mathcal{Q}^{K-M} \\ j' \neq j''}} \delta_i \delta_{j'} \delta_{j''} \frac{\mathbb{C}(m_{i,j'}^{(\mathcal{I})}, m_{i,j''}^{(\mathcal{I})})}{\sigma_{\Omega}^2}}_C \quad (1)$$

The index $\mathbf{T}_{\Omega}^{(\mathcal{I})} = 1 - \mathbf{P}_{\Omega}^{(\sim \mathcal{I})}$ verifies:

$$\mathbf{T}_{\Omega}^{(\mathcal{I})} = \underbrace{\sum_{\substack{i \in \mathcal{Q}^M \\ j \in \mathcal{Q}^{K-M}}} \delta_j \delta_i^2 \frac{\sigma_{i,j}^2}{\sigma_{\Omega}^2} \mathbf{T}_{\omega_{i,j}}^{(\mathcal{I})}}_D + \underbrace{\frac{\sigma_B^2}{\sigma_{\Omega}^2} \tilde{\mathbf{T}}_{\mathcal{Q}^K}^{(\mathcal{I})}}_E - \underbrace{\sum_{\substack{j \in \mathcal{Q}^{K-M} \\ i', i'' \in \mathcal{Q}^M \\ i' \neq i''}} \delta_j \delta_{i'} \delta_{i''} \frac{\mathbb{C}(m_{i',j}^{(\sim \mathcal{I})}, m_{i'',j}^{(\sim \mathcal{I})})}{\sigma_{\Omega}^2}}_F + \underbrace{\sum_{j \in \mathcal{Q}^{K-M}} \delta_j \delta_i (1 - \delta_i) \frac{\sigma_{i,j}^2}{\sigma_{\Omega}^2}}_G \quad (2)$$

- A and D are the contribution of indices computed *within* cuboids,
- B and E are the contribution of indices of the mean of f *between* cuboids,
- C quantifies the similarity of the means of f on the margin $\omega_{\mathcal{I}}$ between the cuboids having the face $\omega_{\mathcal{I}}$ in common,
- F quantifies the similarity of the means of f on the margin $\omega_{\sim \mathcal{I}}$ between the cuboids having the face $\omega_{\sim \mathcal{I}}$ in common,
- G is a constant that depends on variance of factors in cuboids and an expression of their volumes,
- main and total sensitivity indices of a factor $X^{(k)}$ are easily obtained with $\mathcal{I} = \{k\}$.

4 Example

A polynomial test model with five variables is defined on hypercube $[0, 1]^5$. The edges of the hypercube are uniformly cut into Q intervals. That cutting defines Q^5 cuboids. Real sensitivity indices are computed at global and cuboids scale. A nested lhs design is proposed to sample cuboids and points inside them to estimate the terms of the relationships (1) and (2). Different sample size and value of Q are used to evaluate the design.

5 Perspectives

The relationship between sensitivity indices at different scales being established, that work will be continued in the following directions:

- trying to improve the estimation by means of a metamodel,
- finding the relationship between the indices with a hierarchical cutting of the domain,
- determining the relevant scale in environmental studies using a numerical model.

6 References

Saltelli et al (2004) *Sensitivity Analysis*, Wiley, New York.

I.M. Sobol, , S. Kucherenko, Derivative based global sensitivity measures and their link with global sensitivity indices, *Mathematics and Computers in Simulation*, Volume 79, Issue 10, June 2009, Pages 3009–3017.

Acknowledgments

This work was financially supported by the ANR project SimTraces (ANR-2011-CESA-008-01).

Using Gaussian process metamodels for sensitivity analysis of an individual-based model of a pig fattening unit

A. Caldéro^{1,2}, A. Aubry¹, L. Brossard², F. Brun³, J.-Y. Dourmad², Y. Salaün¹, and F. Garcia-Launay²

¹IFIP - Institut du porc, 35651, Le Rheu, France

²PEGASE, INRA, Agrocampus Ouest, 35590, Saint-Gilles, France

³ACTA - le réseau des instituts des filières animales et végétales, 31326, Castanet-Tolosan, France

Abstract

Pig livestock farming systems face economic and environmental issues. To cope with them and identify innovative strategies different models have been developed to predict performance of fattening pig production. However, most of them do not account for the interactions between feeding strategies, management practices, variability of performance and requirements among pigs. Recent studies have highlighted the added value of individual-based models to quantify the effects of feeding practices on technical and environmental performance of a group of pigs (Brossard et al., 2014). Our objective was to develop a pig fattening unit model able i) to simulate individual performance of pigs including their variability in interaction with the farmer's practices and decisions and ii) to evaluate the effects of these practices and decisions on the technical, economic and environmental performance.

The fattening unit model is a discrete-event mechanistic model, stochastic for biological traits (intake and growth potential of pigs, risk of mortality), with a one-day time step. The pig fattening system articulates three subsystems: a biophysical, an operating, and a decision system. The biophysical system contains the fattening pigs. The operating system includes the resources (feeds, fattening rooms), applies the manager's rules to the biophysical system and allocates resources to each activity in the farm. The decision system is represented by the manager (the farmer) who decides, manages, and controls the operating system, and indirectly the biophysical one. The manager receives information from the biophysical system and from the agenda of events each day. According to this information it updates the agenda through the addition/removal of tasks. Practices are inputs of the model, as well as feed composition. The practices include batch management, pen filling practices, feeding practices and slaughter delivery practices. Pigs are represented using an individual-based model adapted from the InraPorc model (van Milgen et al., 2008). This individual-based model simulates the feed intake, body protein and body lipid deposition, and the resulting growth and nutrient excretion of each pig, on a daily basis.

The model is here applied on a typical pig fattening unit in terms of size, batch management, and feeding strategy. The simulation run on 1000 days and performed the dynamic growth of 47 batches (i.e. from 4700 to 42300 pigs). The chosen interval of 21 days between two successive batches corresponds to the main figure in French pig farms. We considered 5 fattening rooms, with the use of an extra-room for the management of too-light pigs. The feed rationing plan simulated the ad libitum distribution of feed. The feed sequence is a ten phase plan mixing two feeds shifting progressively from a 100% grower to a 100% finisher feed. These plans were applied using the mean weight of pigs from the same pen as transition criterion. The range of slaughter body weight for carcass payment without penalties

was between 105 kg to 135 kg, in which the optimal range for carcass payment was between 112 kg to 127 kg, when referring to the French payment grid for lean meat content and carcass weight. For these simulations we controlled the stochasticity by setting the seed in order to have the same scheme of mortality between the simulations. The pig profiles (5 parameters driving growth potential and intake curve) were given as inputs to the model from a dataset containing 1000 female profiles and 1000 male profiles (Brossard et al., 2014).

The model outputs used for the sensitivity analysis are average values for the slaughter age, the slaughter weight, the amount of phosphorus and nitrogen excreted per pig over the fattening period, the lean content per pig, the feed conversion ratio, the total feeding cost per pig, the daily gain, the percentage of pigs in the slaughter weight range, and the one in the optimal range of slaughter weight. Considering the running time for one simulation (around 10 mins), we chose to use Gaussian Process Metamodels to reduce the time cost of the sensitivity analysis. One metamodel was built per output from 100 simulations of the model, using Latin hypercube sampling rescaled with the chosen parameter's distribution. Table 1 shows the 14 parameters tested in the sensitivity analysis and the distribution and bounds associated. The parameters feed, phosphorus, nitrogen and amino acid intake are mainly for checking the behaviour of the model and its correct implementation by expertise. The area per pig and rate of pigs per room aimed to evaluate the sensitivity of the model to density of pigs in pen. The parameters cleaning-disinfection and drying period, size of extra-room and maximum time fatten in extra-room are mainly for evaluate the effect of the duration of the fattening period on the model. The number of place per pen and number of pen per room aimed to study the impact of the size of the farm on the model. The minimum number of pig per delivery to slaughterhouse and tolerance on the number of pig delivered compared to the announcement are mainly to evaluate the effect of the constraint on delivery on the model. The mortality rate aimed to check the effect of this part of the animals' characteristics on the model. The extended Fourier Amplitude Sensitivity Test (eFAST) method (Saltelli et al., 1999) was used to perform the sensitivity analysis on the metamodels, using $N = 1500$ scenarios for each trajectory, which indices 21000 simulations of each metamodel. We performed the sensitivity analysis using the fast99 (Pujolet al., 2015) function in R 3.2.2.

Figure 1 shows the results of the sensitivity analysis. Figure 1.A gives the coefficients of variation of the model's studied outputs. This graph was built using the metamodels output values obtained from the input sequences of the scenario considered above. The greatest variations were observed for the quantity of phosphorus and nitrogen excreted per pig, and the percentage of pigs delivered in the optimal weight of range. These variations of phosphorus and nitrogen excretion are explained respectively at 85% by intake of phosphorus, and at 84 % by the intake of nitrogen. Concerning percentage of pigs in optimal weight range, 38% of the variation is explained by the number of days of the drying period, 18% by the minimum number of pigs required to schedule a delivery, 14% by the quantity of feed intake, and 8% by the number of places per pen. Figure 1.B shows the average sensitivity indices of the 14 inputs investigated among all the outputs. Feed intake explains 37% of the total variation of the model outputs, in particular lean content (88%), average daily gain (88%), feed conversion ratio (59%), feed cost per pig (43%), slaughter weight (25%) and percentage of pigs in optimal weight range (14%). The duration of the drying period explains 31% of the total variation of the model outputs, in particular slaughter age (75%), slaughter weight (52%), feed cost per pig (44%), percentage in optimal weight range (38%) and feed conversion ratio (32%). Phosphorus and nitrogen intake explain each 11% of the total variation of the model outputs. The minimum number of pigs required to schedule a delivery explains 10% of the total variation of the model outputs, in particular percentage of pigs in weight range (34%) and in optimal weight range (18%). The other inputs explain less than 5% of the mean total variation among all the outputs. The maximum of variation explained by the number of places per pen was 16% through the percentage of pigs in weight range.

In conclusion, the sensitivity analysis allowed us to check by expertise that the response of the model corresponds to the results expected, to detect the last informatics errors and correct them, and to

identify the less-sensitivity parameters which can be set for routine use. In this paper we used the first step development of a fattening unit model for the prediction of technical performance. The following step will be to predict the economic results and the environmental impacts at farm gate using LCA. In addition of the sensitivity analysis, an analysis by virtual experiments will be performed on the model.

References

- Brossard L, Vautier B, van Milgen J, Salaün Y and Quiniou N 2014. Comparison of in vivo and in silico growth performance and variability in pigs when applying a feeding strategy designed by simulation to control the variability of slaughter weight. *Animal Production Science* 54, 1939-1945.
- Pujol G, Ioos B and Janon A 2015. The R package “sensitivity”, version 1.11.1. Technical report, CRAN.
- Saltelli A, Tarantola S and Chan KPS 1999. A quantitative model-independent method for global sensitivity analysis of model output. *Technometrics* 41, 39-56.
- van Milgen J, Valancogne A, Dubois S, Dourmad JY, Seve B and Noblet J 2008. InraPorc: A model and decision support tool for the nutrition of growing pigs. *Animal Feed Science and Technology* 143, 387-405.

Table 1. Parameters and associated distribution used for the sensitivity analysis of the model

Parameter	Lower boundary	Upper boundary	Distribution law
Feed intake (Proportion of InraPorc intake)	0.95	1.05	Uniform
Drying period (day)	0	6	Discrete uniform
Phosphorus intake (Proportion of InraPorc intake)	0.9	1.1	Uniform
Nitrogen intake (Proportion of InraPorc intake)	0.9	1.1	Uniform
Min. nb. Pigs/delivery (pig)	10	190	Discrete uniform
Nb. Places/pen (place)	10	30	Discrete uniform
Nb. Pens/room (pen)	10	30	Discrete uniform
Delivery tolerance (Proportion of number of pigs)	0	0.15	Uniform
Size of extra-room (Proportion of farm size)	0	0.15	Uniform
Amino acid intake (Proportion of InraPorc intake)	0.95	1.05	Uniform
Mortality rate (dmnl.)	0.01	0.05	Uniform
Rate pigs/room (Proportion of room size)	0.9	1.1	Uniform
Max. time in extra-room (day)	7	63	Discrete uniform
Area/pig (m ²)	0.6	1	Uniform

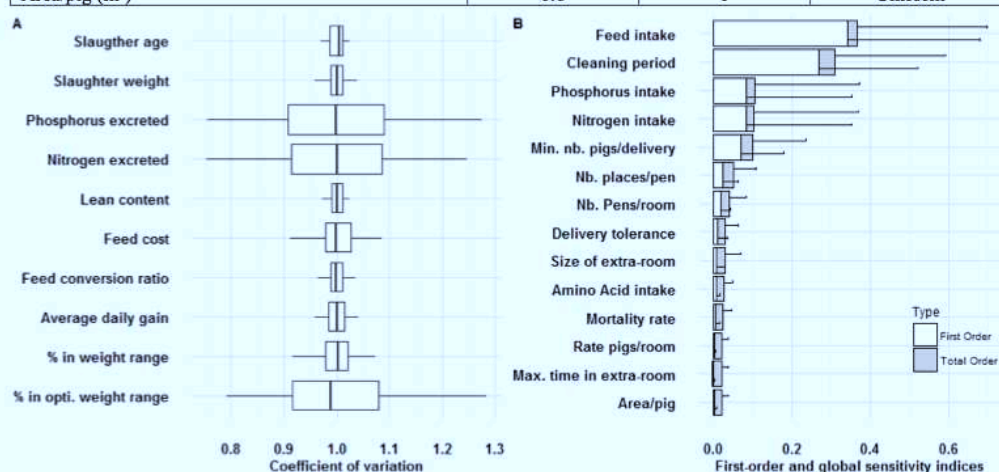


Figure 1. Results of the sensitivity analysis A) coefficients of variation of main outputs of the model and B) average sensitivity indices of the inputs investigated (among all the outputs)

Sensitivity analysis and calibration of a numerical code for the prediction of power from a photovoltaic plant

M. Carmassi^{1,2}, P. Barbillon¹, H. Bouia², M. Keller³, A. Lindsay², and E. Parent¹

¹UMR MIA-Paris, AgroParisTech, INRA, Université Paris-Saclay, Paris, France

²EDF R&D EnerBat Department, Moret-sur-Loing, France

³EDF R&D MRI Department, Chatou, France

Abstract

So far, most studies in the photovoltaic (PV) field have been done with a pseudo- deterministic point of view. The input uncertainty is propagated through a numerical code and provides the results with a 90% interval confidence, the parameters and the input probability distributions being determined by an expert. However, for investors, the risk associated to the investment is closely linked to the uncertain- ties in the evaluation of how much the PV power plant will produce. A statistical framework is needed to provide a more accurate estimation of the power produced by the PV plant and the associated error. For the modeling of PV power plants, a variety of computer models have been constructed. One of them, which is also the most accurate to date, is highly time-consuming. Hence, this study will deal with the sensitivity analysis and calibration of time-consuming codes, based on a large amount of data.

The modeling of a PV power plant can be performed with an equivalent electric scheme which tries to match its physical behavior. Two physical models are introduced. The first one is "simple" and the physical equations are straightforward. A Python code can be made up from those equations and predicts the amount of power, generated by a number of panels in simplified environmental conditions, relatively quickly. The second one is more complete. It matches better the real behavior of the PV power plant and can account for the partial shadings that may occur in a large-size power plant. Thus, the second physical model is more interesting to work with but also more complicated (high computing time and less regular behavior for example).

A physical model has two kinds of inputs: controlled variables which are observed in experimental conditions on the one hand and (usually uncertain) parameters on the other hand. The latter have to be calibrated to make the outputs of the physical model close to the observed quantities of interest, for instance the instantaneous power of the PV plant. Controlled variables are, for example, the meteorological data, consisting of: the amount of irradiation from the sun, the temperature, the geographical position (latitude and longitude) and the time. The parameters are factors inherent to the physical model (the yield of the PV module, the module temperature coefficient, etc.). Generally, the values of these factors are fixed according to expert elicitation. However, real-life experience shows that, when parameters are set according to expert opinion only, code outputs may be far from experimental data. To determine and evaluate the uncertainties on these parameters more precisely, a calibration has to be conducted. A Bayesian framework is adopted to make this inference. This will allow us to confront expert information and experimental data.

The first step for such a study is to perform a sensitivity analysis. This is crucial for the following step because sorting parameters from the most to the less important will allow us to save a lot of

computational time thereafter. The number of parameters to be calibrated depends on the physical model. However in both present cases it always exceeds ten. A screening method is first carried out to separate the ones which have no overall impact on the output. Afterwards, a Sobol analysis is done to sort the remaining parameters and indicate which one is the most important. The Sobol analysis is important in this case, because it provides a physical point of view on which parameters have an impact on the power and allows us to check the accuracy of the analysis.

The second step is to calibrate the parameters. Once a statistical model is set, Bayesian inference will combine all available data with a prior distribution to obtain a posterior distribution on the unknown parameters. The experimental data are the power measured instantaneously on the test stand. At a high sampling frequency, the number of experimental points is very high. However, all of them are not really informative. For example, during the night the PV power plant produces nothing and these data can be removed from the data set. Globally, the considered data can be limited to the time period from daybreak to sunset. The calibration is performed with Markov Chains Monte Carlo algorithms (MCMC) like Metropolis- Hastings or Metropolis within Gibbs algorithms. A tempering scheme is adopted to deal with the full amount of available data without jeopardizing the time needed to achieve convergence of MCMC algorithms. Furthermore, an adaptive exploration kernel is chosen for the MCMC algorithms since the dimension is high and adapting the exploration kernel to the covariance of the posterior distribution will accelerate convergence.

To ensure that the production predictions of the power plant provided by the code are reliable, the code has to be validated. Validation means to assess whether the code produces outputs close to observed power measures once calibration has been conducted. This validation question can be expressed as a choice between two statistical models. In the first one, the only error between the outputs of the physical model and the observed power measures is a measurement error i.e. a classical white noise process. In the second one, a discrepancy term will be added to the measurement error to capture a systematic error of the physical model. Usually, this discrepancy is modeled as a Gaussian process. If the model with the discrepancy is selected by the statistical procedure, it will mean that the physical model does not predict well enough the actual power. In this case, pure model predictions can be corrected by adding the estimated discrepancy.

Finally, all of these techniques are computationally demanding. Indeed, some of the physical models are expensive and a Gaussian process emulator has been introduced to reduce the overall computation load.

A New Bayesian Approach for Statistical Calibration of Computer Model

Frédéric Delay^{*1}, Thierry Mara², and Anis Younes^{1,3}

¹Laboratoire Hydrologie et Geochimie de Strasbourg, University of Strasbourg/EOST, CNRS, 1 rue Blessig
67084 Strasbourg, France

²PIMENT, Université de La Réunion, 15 Avenue René Cassin, BP 7151, 97715 Moufia, La Réunion, France

³IRD UMR LISAH, F-92761 Montpellier, France

Abstract

The validation of computer models is an essential task to increase their credibility. One of the most important exercises in the validation framework is to check whether the computer model adequately represents reality [1]. This is achieved by comparing model predictions to observation data. This exercise generally leads to model calibration because the model parameters are usually poorly known a priori (i.e. before collecting data). Good practice in calibration of computer models consists of searching for all parameter values that satisfactorily fit the data, thus determining their plausible range of uncertainty. This can be achieved in a Bayesian framework in which the prior knowledge about the model and the observed data are merged to define the joint posterior probability distribution function (pdf) of the parameters. The issue is then to assess the joint posterior pdf.

In probabilistic inverse modeling, the parameter set $\mathbf{x} = (x_1, \dots, x_d)$ of a computer model is inferred from a set of observation data $\mathbf{y}$ using the Bayesian inference, which defines the conditional joint posterior pdf as follows:

$$p(\mathbf{x}|\mathbf{y}) \propto p(\mathbf{y}|\mathbf{x})p(\mathbf{x}), \quad (1)$$

where $p(\mathbf{x})$ is the prior density that characterizes the investigator's beliefs about the parameters before collecting the new observations, and $p(\mathbf{y}|\mathbf{x})$ is the likelihood function, which measures how well the model fits the data. The parameter set that maximizes Eq. (1), namely:

$$\mathbf{x}^{\text{MAP}} = \underset{\mathbf{x}}{\operatorname{argmax}} p(\mathbf{x}|\mathbf{y}), \quad (2)$$

is called the Maximum A Posteriori (MAP) estimate of the parameters. It is the most probable parameter set given the data and can be inferred via an optimization technique. The marginal posterior pdf that characterizes the uncertainty of a single parameter is defined by the following integral:

$$p(x_i|\mathbf{y}) = \int p(\mathbf{x}|\mathbf{y}) d\mathbf{x}_{-i}, \quad \forall i = 1, \dots, d \quad (3)$$

where $\mathbf{x}_{-i}$ represents all the parameters except x_i . Usually, the integral in Eq. (3) is evaluated by a multidimensional quadrature method or by direct summations in a large sample of $p(x_i|\mathbf{y})$ obtained, for instance, via an MCMC technique.

^{*}Corresponding author: fdelay@unistra.fr

The inference of model parameter posterior pdf by means of Markov chain Monte Carlo (MCMC) sampling techniques [2, 3] has received much attention in the last two decades. MCMC explores the region of plausible values in the parameter space and provides successive parameter draws directly sampled from the target joint pdf. Some selection criteria are used to ensure that the successive draws in the chain improve. This means that, throughout the sampling process, probable draws with respect to the target distribution are more likely drawn. Many developments and improvements have been proposed to accelerate MCMC convergence (see [4–6]).

Recently, [7] proposed a new probabilistic approach to the inverse problem whose main idea is to maximize the joint posterior pdf of a parameter set with one selected parameter sampling successive prescribed values. This provides the so-called Maximal Conditional Posterior Distribution (MCPD) of the selected parameter. The maximal conditional posterior distribution of x_i is defined as follows:

$$\mathcal{P}(x_i) = \max_{\mathbf{x}_{-i}} (p(\mathbf{x}_{-i}|\mathbf{y}, x_i)) \times p(x_i|\mathbf{y}). \quad (4)$$

An informal definition can be given by stating that a point estimate of the MCPD is the maximal value reached by the joint pdf (Eq. (1)) for a given (prescribed) value of one parameter i.e. x_i . This maximal value, in the context of model inversion, assumes that the set $\mathbf{x}_{-i}$ maximizes (Eq. (1)), knowing that x_i is prescribed. By applying the axiom of conditional probabilities to (Eq. (4)), it can be stated that $\max \{p(\mathbf{x}_{-i}|\mathbf{y}, x_i)\} \times p(x_i|\mathbf{y}) = \max_{\mathbf{x}_{-i}} \{p(\mathbf{x}_{-i}, x_i|\mathbf{y})\}$. Therefore, the MAP estimate (when it exists) belongs to the MCPD of all parameters.

The main advantage of the recent MCPD technique is that parameter distributions can be inferred independently. Therefore, the MCPDs can be simultaneously evaluated on multicore computers (or on multiple computers). This drastically reduces the computational effort in terms of computational time units (CTU). Our presentation will develop the MCPD approach and exemplify its efficiency in solving inverse problems.

References

- [1] M. J. Bayarri, J. O. Berger, R. Paulo, J. Sacks, J. A. Cafeo, J. Cavendish, C. H. Lin, and J. Tu. A framework for validation of computer models. *Technometrics*, 49:138–154, 2007.
- [2] N.-A. Metropolis, A. W. Rosenbluth, M. N. Rosenbluth, A. H. Teller, and E. Teller. Equations of state calculations by fast computing machines. *Journal of Chemical Physics*, 21:1087–1091, 1953.
- [3] H. Hastings. Monte Carlo sampling methods using Markov chains and their applications. *Biometrika*, 57:97–109, 1970.
- [4] H. Haario, M. Laine, A. Mira, and E. Saksman. DRAM: Efficient adaptive MCMC. *Statistics and Computing*, 16(339–354), 2006.
- [5] C. J. F. ter Braak and J. Vrugt. Differential evolution markov chain with snooker updater and fewer chains. *Statistics and Computing*, 18(4):435–446, 2008.
- [6] E. Laloy and J. A. Vrugt. High-dimensional posterior exploration of hydrologic models using multiple-try DREAM_(zs) and high-performance computing. *Water Resources Research*, 48:W01526, 2012.
- [7] T. A. Mara, N. Fajraoui, A. Younes, and F. Delay. Inversion and uncertainty of highly parameterized models in a bayesian framework by sampling the maximal conditional posterior distribution of parameters. *Advances in Water Resources*, 76:1–10, 2015.

Abstract submission for SAMO 2016 conference - Adaptive numerical designs for the calibration of computer codes

Guillaume Damblin*

CEA SACLAY DEN/DANS/DM2S/STMF/LGLS, F-91191 Gif-sur-Yvette Cedex

April 22, 2016

Abstract

After having performed a sensibility analysis for screening influential inputs of a computer code, practitioners should aim at making the code outputs as close as possible to a set of field experiments in order to improve its predictive capability. That issue is called *calibration* (Campbell, 2006).

Our framework deals with a scalar physical quantity of interest, referred to as $r(\mathbf{x})$, where $\mathbf{x}$ is a vector of control variables and with a computer code $y_{\boldsymbol{\theta}}(\mathbf{x})$ where $\boldsymbol{\theta} \in \mathcal{T}$ is a vector of parameters having no observable counterpart in the reality and thus most often uncertain. The goal of statistical calibration consists in reducing the uncertainty affecting $\boldsymbol{\theta}$ with the help of a statistical model which links the code outputs with the field measurements, denoted by $\mathbf{z}^f := (z_1^f, \dots, z_n^f)$ which are related to n experimental sites $\mathbf{X}^f = [\mathbf{x}_1^f, \dots, \mathbf{x}_n^f]^T$. By assuming no code discrepancy can occur between $y_{\boldsymbol{\theta}}(\mathbf{x})$ and $r(\mathbf{x})$ for any $\mathbf{x} \in \mathcal{X}$, we have for $1 \leq i \leq n$:

$$z_i^f = y_{\boldsymbol{\theta}}(\mathbf{x}_i^f) + \epsilon_i, \quad (1)$$

where

$$\epsilon_i \sim \mathcal{E}_i \underset{i.i.d.}{=} \mathcal{N}(0, \lambda^2)$$

statistically encodes both the residual variability and the measurements error of the physical experiment (Kennedy and O'Hagan, 2001). In a Bayesian setting, where λ^2 is assumed known, the posterior distribution of $\boldsymbol{\theta}$ is then written as

$$\begin{aligned} \pi(\boldsymbol{\theta}|\mathbf{z}^f) &\propto \mathcal{L}(\mathbf{z}^f|\boldsymbol{\theta})\pi(\boldsymbol{\theta}), \\ &\propto \frac{1}{(\sqrt{2\pi}\lambda)^n} \exp\left[-\frac{1}{2\lambda^2}SS(\boldsymbol{\theta})\right]\pi(\boldsymbol{\theta}), \end{aligned} \quad (2)$$

where

$$SS(\boldsymbol{\theta}) = \|\mathbf{z}^f - y_{\boldsymbol{\theta}}(\mathbf{X}^f)\|^2 \quad (3)$$

is the sum of squares of the residuals between the simulations and the field measurements. It is usually sampled using MCMC methods that become infeasible when the simulations are highly time-consuming. A way to circumvent this issue consists in replacing the computer code with a Gaussian Process Emulator (GPE) (Santner et al., 2003). It is built thanks

*Electronic address: guillaume.damblin@cea.fr

to a learning sample of simulations $y(\mathbf{D}_M)$ run over a design of experiments $\mathbf{D}_M$. Then, a surrogate posterior distribution π^S based on the GPE can be established:

$$\pi^S(\boldsymbol{\theta}|\mathbf{z}^f, y(\mathbf{D}_M)) \propto \mathcal{L}^S(\mathbf{z}^f|y(\mathbf{D}_M), \boldsymbol{\theta})\pi(\boldsymbol{\theta}). \quad (4)$$

where

$$\mathcal{L}^S(\mathbf{z}^f|y(\mathbf{D}_M), \boldsymbol{\theta}) \propto |V_{\hat{\boldsymbol{\Psi}}, \hat{\sigma}^2}^M(\boldsymbol{\theta}) + \lambda^2 \mathbf{I}_n|^{-1/2} \exp \left\{ -\frac{1}{2} \left[(\mathbf{z}^f - \mu_{\hat{\boldsymbol{\beta}}, \hat{\boldsymbol{\Psi}}}^M(\mathbf{D}_{\boldsymbol{\theta}}))^T (V_{\hat{\boldsymbol{\Psi}}, \hat{\sigma}^2}^M(\boldsymbol{\theta}) + \lambda^2 \mathbf{I}_n)^{-1} (\mathbf{z}^f - \mu_{\hat{\boldsymbol{\beta}}, \hat{\boldsymbol{\Psi}}}^M(\mathbf{D}_{\boldsymbol{\theta}})) \right] \right\} \quad (5)$$

is the conditionnal likelihood of $\mathbf{z}^f$ with respect to $y(\mathbf{D}_M)$ where $(\hat{\boldsymbol{\beta}}, \hat{\boldsymbol{\Psi}}, \hat{\sigma}^2)$ are plug-in estimators of the GPE's parameters. The surrogate posterior (4) and the target posterior (2) are different in that $y_{\boldsymbol{\theta}}(\mathbf{X}^f)$ is replaced by the mean vector of the GPE $\mu_{\hat{\boldsymbol{\beta}}}^M(\mathbf{D}_{\boldsymbol{\theta}})$ and the conditional covariance matrix $V_{\hat{\boldsymbol{\Psi}}, \hat{\sigma}^2}^M(\boldsymbol{\theta})$ is added up to $\lambda^2 \mathbf{I}_n$.

By doing so, the surrogate posterior can be sampled using MCMC methods instead of the target one, but this is subject to an error which strongly depends on the numerical design of experiments $\mathbf{D}_M$ used to fit the GPE. The most used default strategy consists in building a Space Filling Design (SFD), such as an optimized Latin Hypercube (Morris and Mitchell, 1995). Our numerical tests have actually shown that they do not work well, leading to large errors in terms of the Kullback Leibler (KL) divergence (Cover and Thomas, 1991) between the surrogate posterior and the target posterior, that is written:

$$\text{KL}(\pi(\boldsymbol{\theta}|\mathbf{z}^f) || \pi^S(\boldsymbol{\theta}|\mathbf{z}^f, y(\mathbf{D}_M))) = \int_{\mathcal{T}} \pi(\boldsymbol{\theta}|\mathbf{z}^f) \left(\log(\pi(\boldsymbol{\theta}|\mathbf{z}^f)) - \log(\pi^S(\boldsymbol{\theta}|\mathbf{z}^f, y(\mathbf{D}_M))) \right) d\boldsymbol{\theta}. \quad (6)$$

Instead of using SFD, we propose to build in an adaptive fashion a proper design limited to $\mathbf{X}^f \times \mathcal{T}$ by means of the Expected Improvement criterion (Jones et al., 1998). Our simulation studies performed on several toy functions in 2d and 6d have shown the efficiency of the sequential strategies for reducing the KL divergence (6).

References

- Campbell, K. (2006). Statistical calibration of computer simulations. *Reliability Engineering and System Safety*, 91(10-11):1358–1363.
- Cover, T. and Thomas, J. (1991). *Elements of Information Theory*. Wiley-Interscience.
- Jones, D., Schonlau, M., and Welch, W. (1998). Efficient global optimization of expensive black-box functions. *Journal of Global Optimization*, 13(4):455–492.
- Kennedy, M. and O'Hagan, A. (2001). Bayesian calibration of computer models (with discussion). *Journal of the Royal Statistical Society, Series B, Methodological*, 63(3):425–464.
- Morris, D. and Mitchell, J. (1995). Exploratory designs for computational experiments. *Journal of Statistical Planning and Inference*, 43:381–402.
- Santner, T., Williams, B., and Notz, W. (2003). *The Design and Analysis of Computer Experiments*. Springer-Verlag.

Global sensitivity analysis in dynamic MFA

Nada Dzubur¹, Hanno Buchner¹, and David Laner¹

¹Institute for Water Quality, Resource and Waste Management, Vienna University of Technology, Karlsplatz 13, 1040 Vienna, Austria

Abstract

Material flow analysis (MFA) is a tool to quantify the flows and stocks of materials in arbitrarily complex systems. Dynamic MFA is a frequently used method to assess past, present and future stocks and flows of materials in the anthroposphere (Müller et al. 2014). In contrast to static MFA, where material flows are determined for one balancing period and therefore time independent, material stocks and flows in a dynamic material flow model can potentially depend on all previous states of the system (Baccini and Bader 1996). Recently, dynamic MFA has become increasingly popular with the primary focus on the investigation of material stocks in society and associated end-of-life flows (cf. Laner and Rechberger, in press). Since models represent a simplification of the real metabolic system and because of data limitations in terms of quality and quantity, uncertainty is inherent to material flow analysis (MFA) (Laner et al. 2014). Therefore, uncertainty is a basic aspect of material flow modelling and needs to be explicitly considered to reduce uncertainties and inconsistencies as far as possible, thereby allowing for reliable decision support (Gottschalk et al. 2010, Laner et al. 2015). With respect to dynamic MFAs, the in-use stocks and end-of-life (EOL) material flows are typically estimated according to a top-down approach (i.e. accounting of the net flows into or out of the stock over time), where substantial uncertainty exists concerning model parameters such as average product lifetimes or historical material use patterns.

In order to understand the effect of limited data quality and model assumptions on MFA results, the use of sensitivity analysis methods in dynamic MFA studies has been on the increase. So far, the usual sensitivity analysis in dynamic MFA is a One-at-a-time method, which is testing parameter perturbations individually and observing the outcomes on output. In contrast to that, variance based global sensitivity analysis decomposes the variance of the model output into fractions caused by the uncertainty or variability of input parameters (Saltelli et al. 2008). The process of recalculating outcomes under alternative assumptions to determine the impact of variables using global sensitivity analysis can be useful to identify model inputs that cause significant uncertainty in the output in order to increase robustness of the model and understanding of the relationships between input and output variables (Panell 1997). Interaction and time-delayed effects of uncertain parameters on the output of an archetypal input-driven dynamic material flow model using a sample based approach for variance based global sensitivity analysis proposed by Saltelli et al. (2008) are investigated in this study. The results show that determining the main (or first-order) effects of parameter variations is often sufficient in dynamic MFA, because substantial effects due to the simultaneous variation of several parameters (higher-order effects) do not appear for classical set ups of dynamic material flow models.

Higher order effects may be relevant for secondary raw material production flows considering sorting and upgrading processes in advance because the probability density function for the respective sector split is located close to zero and several other parameters are multiplied with the sector split ratio

to calculate the flow of interest. For models with time-varying parameters, time delay effects of parameter variation on model outputs need to be considered, potentially boosting the computational cost of global sensitivity analysis. The implications of exploring the sensitivities of model outputs with respect to parameter variations in the archetypical model are used to derive model- and goal-specific recommendations on choosing appropriate sensitivity analysis methods in dynamic MFA. Dynamic material flow models will gain in complexity in the future due to the consideration of various material quality layers (e.g. Buchner et al. 2015) or the requirement of closed mass balances applied to the model (i.e. recycled material flows have to (exactly) correspond with the quantities used in secondary production). Because higher order effects are expected to become more prominent in such models, the investigation of parameter interaction effects and parameter dependencies (e.g. Mara et al. 2015) will become a major field for extending the use of sensitivity analysis in dynamic MFA.

References

- Baccini P, Bader HP. *Regionaler Stoffhaushalt: Erfassung, Bewertung und Steuerung*: Spektrum, Akad. Verlag; 1996.
- Buchner H, Laner D, Rechberger H, Fellner J. Future Raw Material Supply: Opportunities and Limits of Aluminium Recycling in Austria. *Journal of Sustainable Metallurgy*. 2015;1(4):253-62.
- Gottschalk, F., R. W. Scholz, and B. Nowack. 2010. Probabilistic material flow modeling for assessing the environmental exposure to compounds: Methodology and an application to engineered nano- TiO₂ particles. *Environmental Modelling & Software* 25(3): 320-332.
- Laner D, Rechberger H. Material flow analysis. Chapter 7 “Special Types of Life Cycle Assessment” (Finkbeiner M ed): *LCA Compendium - The Complete World of Life Cycle Assessment* (Kloppfer W, Curran MA, series eds). Springer, Dordrecht; 2016.
- Laner, D., H. Rechberger, and T. Astrup. 2014. Systematic Evaluation of Uncertainty in Material Flow Analysis. *Journal of Industrial Ecology* 18(6): 859-870.
- Laner, D., H. Rechberger, and T. Astrup. 2015. Applying Fuzzy and Probabilistic Uncertainty Concepts to the Material Flow Analysis of Palladium in Austria. *Journal of Industrial Ecology*: n/a-n/a.
- Mara T.A., Tarantola S., Annoni P. Non-parametric methods for global sensitivity analysis of model output with dependent inputs. *Environmental Modelling & Software*. 2015;72:173-83.
- Müller E, Hilty LM, Widmer R, Schluep M, Faulstich M. Modeling Metal Stocks and Flows: A Review of Dynamic Material Flow Analysis Methods. *Environmental Science & Technology*. 2014;48(4):2102-13.
- Pannell D.J. *Introduction to practical linear programming*: J. Wiley; 1997.
- Saltelli A., Ratto M., Andres T., Campolongo F., Cariboni J., Gatelli D., et al. *Global Sensitivity Analysis: The Primer*: Wiley; 2008.

Development of a High Performance Capabilities for Supporting Spatially-Explicit Uncertainty- and Sensitivity Analysis in Multi-Criteria Decision Making

Christoph Erlacher¹, Piotr Jankowski², Seda Şalap-Ayça², Karl-Heinrich Anders¹, Gernot Paulus¹

¹ Carinthia University of Applied Sciences – Department of Geoinformation and Environmental Technologies, Austria

² San Diego State University – Department of Geography, United States of America

Spatial multi-criteria decision making (MCDM) applications often do not provide any detailed information about the robustness of results and uncertainties associated with input data. Applications from landscape assessment, natural hazard risk assessment for communities, identification of land use strategies for sustainable regional development, water resource management or habitat suitability in the context of environmental protection are just a few examples of domain areas where MCDM methodology continues to find use.

Therefore, the main objective of this exploratory research project is the development of scalable and adaptable capabilities to accelerate Spatially-Explicit Uncertainty and Sensitivity Analysis (SEUSA). A parallel algorithm design for the implementation of the SEUSA framework will allow reasonable computational times making this kind of spatial analysis applicable and attractive in the context of MCDM.

A spatial MCDM approach considered in this research involves a certain number of alternatives, which are evaluated on the basis of conflicting criteria. Criteria can be represented either as factor- or constraint maps (limitations). In general, a MCDM-based workflow involves the standardization of the criteria to achieve their comparability, expert preferences (weights) that represent the influence of the criteria, and mathematical functions (Weighted Linear Combination, Ideal Point, Ordered Weighted Averaging, Analytic Hierarchy Process etc.) to generate the suitability surface. Detailed information about the spatial MCDM workflow can be found in Malczewski (1999), Erlacher et al. (2009) and Malczewski and Rinner (2015).

One important part of this workflow is the sensitivity analysis to validate the robustness and stability of implemented MCDM models. Uncertainties can be caused by imprecise data or measurement errors, the standardization of criteria values, the implemented decisions rules, and the preferences of the experts expressed by cardinal weights. Uncertainty analysis quantifies the variability of model outcomes and sensitivity analysis focuses on identifying decision criteria or criteria weights that cause the variability (Ligmann-Zielinska and Jankowski, 2014). Spatially-explicit U-SA as research topic was illuminated and discussed in a small number of contributions (Ligmann-Zielinska and Jankowski, 2012; Feizizadeh et al., 2014; Ligman-Zielinska and Jankowski, 2014; Şalap-Ayça and Jankowski, 2016). These studies refer to a variance-based U-SA approach that includes the quasi-random Sobol's experimental design for generating the weight samples and the Monte Carlo Simulation (MSC) to create the suitability surfaces according to the implemented decision rules. This approach explicitly accounts for the interaction of the

input factors, which can be a time-consuming process especially in case of spatial data. The computational effort depends on the number of criteria, the number of raster cells (pixels) and the number of simulations.

First intermediate results were presented at the AAG conference 2016 in San Francisco and focused on the development of a GPU-based (Graphics Processing Units) prototype to accelerate the spatially-explicit U-SA. The implemented prototype refers to a land-allocation problem in order to prioritize agricultural land units based on environmental benefits (Şalap-Ayça and Jankowski, 2016). Details regarding the conceptual development and the CUDA implementations are described in Erlacher et al. (2016).

For the development of a scalable and adaptable approach to accelerate spatially-explicit U-SA several projects involving diverse application areas have to be investigated, in order to support different kinds of MCDM models and operations (local-, focal-, zonal and global raster- and vector analysis) as well as to identify limitations. The successful realization of the project will have a positive impact on future uses of MCDM methodology due to computational support for spatially-explicit uncertainty and sensitivity analysis.

In summary, the proposed research will support spatial decision support capabilities through increased traceability, objectivity, and transparency of results obtained from applications of MCDM.

References

- Erlacher, C. M., Paulus, G. & P. Bachhiesl; (2009). Comparison and Analysis of Multiple Scenarios for Network Infrastructure Planning. In: Car, A., Griesebner, G. & J. Strobl (eds.): *Geospatial Crossroads @ GI_Forum'09*, Wichmann, Berlin/Offenbach, pp.26-35.
- Erlacher, C., Şalap-Ayça, S., Jankowski, P., Anders, K., H., Paulus, G. (2016). A GPU-based Solution for Accelerating Spatially-Explicit Uncertainty- and Sensitivity Analysis in Multi-Criteria Decision Making. *Spatial Accuracy 2016*, in press.
- Feizizadeh, B., Jankowski, P., Blaschke, T. (2014). A GIS based Spatially-explicit Sensitivity and Uncertainty Analysis Approach for Multi-Criteria Decision Analysis. *Computers and Geosciences*, Volume 64, pp. 81-95.
- Ligmann-Zielinska, A., Jankowski, P. (2014). Spatially-Explicit Integrated Uncertainty and Sensitivity Analysis of Criteria Weights in Multicriteria Land Suitability Evaluation. *Environmental Modelling & Software*, Volume 57, pp. 235-247.
- Ligmann-Zielinska, A., Jankowski, P. (2012). Impact of proximity-adjusted preferences on rank-order stability in geographical multicriteria decision analysis. *J. Geogr. Syst.* Volume 14, Issue 2, pp. 167-187.
- Malczewski, J. (1999). *GIS and Multicriteria Decision Analysis*. New York, Wiley.
- Malczewski, J. (2011). Local Weighted Linear Combination. *Transactions in GIS*, Volume 15, Issue 4, pp. 439-455.
- Malczewski, J., Rinner, C. (2015). *Multicriteria Decision Analysis in Geographic Information Science, Advances in Geographic Information Science*. New York, Springer.
- Şalap-Ayça, S., Jankowski, P. (2016). Integrating Local Multi-Criteria Evaluation with Spatially Explicit Uncertainty-Sensitivity Analysis, *Spatial Cognition & Computation*, DOI: 10.1080/13875868.2015.1137578

Sensitivity and uncertainty analyses in climate research: Tool development and application for an atmosphere model

M. Flechsig¹, S. Molnos¹, and T. Nocke¹

¹Potsdam Institute for Climate Impact Research, 14473 Potsdam, Germany

Abstract

In climate and climate impact research model development and application are important methodologies. Numerical models of the Earth system and its subsystems play a key role in understanding the physical processes and in assessing implications of future climate change. Typically, such models are characterized by a high complexity of nonlinear processes with threshold effects and strong interactions, an intrinsic variability of processes, a large number of uncertain model factors, and large volume of multi-variate and multi-dimensional output. In recent years, there has been a growing demand for complementing the findings from simulation experiments with sensitivity and uncertainty measures and to share such information with the scientific community as well as with policy and decision makers (e.g., Katz et al. 2013 and IPCC 2014).

Typically, such simulation models are implemented in programming languages rather than modelling systems. Beside validation tasks, simulation experiments are mainly performed for scenario and re-analyses. For simulation studies with the focus on sensitivity and uncertainty analyses (SUA) methodical challenges arise from interfacing the model to appropriate tools, experiments for sensitivity and uncertainty analyses (SUA) in high-dimensional factor spaces, load distribution of the run ensemble, specification of experiment-specific measures during experiment analysis, and visual analytics of experiment output and derived SUA measures.

SimEnv (Flechsig et al. 2013) presented in this paper, is a multi-run simulation environment for SUA of multi input / output models that meets most of the above criteria: Experiment design is based on pre-defined deterministic, probabilistic and Bayesian experiment types for factor spaces of any dimension that only have to be equipped with numerical information. Experiments cover variance based and Monte Carlo techniques, local and qualitative sensitivity analysis, (fractional) factorial designs, Bayesian calibration, and one-criterial optimization.

The simulation environment comes with a simple model interface that requires only minimal source code modifications of C/C++, Fortran, Java, Python, Matlab, Mathematica, GAMS or shell script models. Multi-variate / -dimensional experiment output is stored in self-describing NetCDF data format. The environment allows for flexible load distribution strategies of the single runs from the run ensemble, supporting multi-core processor machines and compute clusters. In experiment analysis, chains of built-in and user-defined operators are applied to multi-dimensional experiment output over the factor space, to external (reference) data, and to other SimEnv experiments to derive SUA measures from secondary experiment output.

SimEnv is coupled to the visualization system SimEnvVis (Nocke 2007) for interactive explorative visual data analysis. It exploits metadata from experiment design and analysis to select appropriate visualization techniques. One of the advantages of SimEnvVis is the ability to cope with multi-run datasets by special visualization techniques like parallel coordinates and graphical tables.

SimEnv has been used for modelling studies with different objectives (e.g., Knopf et al. 2006 and 2008, and van Oijen et al. 2013). Here, we applied it to study the atmosphere model Aeolus (Coumou et al. 2011). Aeolus is a statistical-dynamical atmospheric model based on time-averaged equations, and therefore much faster than the more widely used atmospheric general circulations models. With Aeolus it is possible to run climate simulations up to multi-millennia timescales and its computational efficiency enables experiment settings with high computational costs in terms of number of runs.

In a first step, we applied simulated annealing technique to optimally tune the models representation of the Hadley cells as well as wind velocities to available observational data. Next, we studied the sensitivity of large-scale atmospheric circulation patterns to different forcing patterns. In particular, we were interested in quantifying the sensitivity of the Hadley circulation to the key dynamical forcings involved and the likely causes behind the observed strengthening and widening of the Hadley cell in recent decades. With appropriate methods we investigated the impact of additional parameters to the Hadley cell's strength and position. Afterwards, we identified and examined in a two-level approach the most sensitive parameters of Aeolus to the Hadley cell's dynamics. For all settings we applied visual data analysis in the coupled multi-dimensional factor - state space.

References

- Coumou, D., Petoukhov, V. and Eliseev, A. V.: Three-dimensional parameterizations of the synoptic scale kinetic energy and momentum flux in the Earth's atmosphere, *Nonlin. Processes Geophys.*, 18(6), 807-827, 2011.
- Flechsig, M., Böhm, U., Nocke, T., Rachimow, C.: The multi-run simulation environment SimEnv. User's Guide, Potsdam Institute for Climate Impact Research, Potsdam, Germany, 2013, <http://www.pik-potsdam.de/software/simenv/>
- IPCC: Climate Change 2014: Synthesis Report. Summary for Policy-makers. IPCC, Geneva, Switzerland, 31 pp., 2014.
- Katz, R.W., Craigmile, P.F., Guttorp, P., Haran, M., Sansó, B., Stein, M.L.: Uncertainty analysis in climate change assessments, *Nature Climate Change*, 3, 769-771, 2013.
- Knopf, B., Flechsig, M., Zickfeld, K.: Multi-parameter uncertainty analysis of a bifurcation point *Nonlin. Processes Geophys.*, 13:531-540, 2006.
- Knopf, B., Zickfeld, K., Flechsig, M., Petoukhov, V.: Sensitivity of the Indian monsoon to human activities *Advances in Atmospheric Sciences*, 25/6:932-945, 2008.
- Nocke, T.: Visual data mining and visualization design for climate research, PhD Thesis, University of Rostock (in German), 2007.
- van Oijen, M. et al.: Bayesian calibration, comparison and averaging of six forest models, using data from Scots pine stands across Europe, *Forest Ecology and Management*, 289:255-268, 2013.

Using the sensitivity analysis to optimize passive cooling solutions in the urban tropical environment

A. Foucquier¹, M. Boulinguez², A. Jay¹, and E. Wurtz¹

¹CEA

²Integrale

Abstract

In France, the building sector is considered as the first energy consumer with at least 40% of the total energy consumption. Many different parameters influence the building behavior as the climate, the envelope characteristics or else the occupancy. In metropolitan France, the heating load constitutes the main point of interest and many solutions have been proposed recently in order to reduce drastically the heating consumption. Inversely, the French tropical territories as the Reunion Island do not know such kinds of heating issues but are more constrained by their huge cooling loads. Indeed, the building sector in tropical environment is affected by energy waste due mainly to the importance of cooling systems consumption.

The sensitivity analysis is increasingly used in the building field since a few years. Initially, the main objectives were to explain better the building behavior and to understand the uncertainties of the numerical building models. In the last few years, the applications of the sensitivity analysis in the building field have been largely diversified. Among them, we will focus in this paper on two specific numerical applications. First, the sensitivity analysis is used as a multi-objective optimizer in order to improve the design of the building envelop considering several different criteria as the thermal comfort and the natural lighting. Secondly, the sensitivity analysis is used to set up efficient monitoring strategies with the aim to improve the thermal comfort during the building operation.

As part of this numerical study, a retrofitted office building will be our case of study. It is based in the coastal town of Saint-Pierre located in the South of the Reunion Island. As we mentioned above, the main issue in the tropical climate concerns the reduction of the cooling consumption. Moreover, this aspect is even more accentuated in office buildings where the internal loads due to computer hardware can become really important. Given that, cooling solutions must be expected to maintain the indoor thermal comfort. First of all, an important aspect occurring during the early design phase and conditioning the future operation phase is the building envelope design defined to anticipate the discomfort caused by the climate and environment local constraints. The Reunion Island knows important solar radiations that can degrade the indoor thermal comfort or increase the cooling consumption. A first step will then be to optimize the building envelope in order to reduce the solar internal loads. However, minimizing the solar gains can lead to deteriorate the visual comfort. Thus, the resulting building envelope must be able to reach the objectives of both thermal and visual comfort. Secondly, in anticipation of the future building operation, other transient cooling solutions able to evacuate instantaneously the exceeded internal loads must be added in order to maintain a suitable thermal comfort. Air conditioning has been largely used in the last few years resulting in the huge increase of energy consumption in the tropical climate. However, some passive cooling solutions have already shown their efficiency and particularly the natural ventilation. The Reunion Island is greatly swept by winds with the trade winds during the day and the thermal breezes during the night. Thus, combined with an efficient monitoring strategy, the potential of natural ventilation

in the Reunion Island can be really effective. Though, implementing monitoring strategies of natural ventilation can be quite complicated by the fact that the natural ventilation depends on many different parameters whereas the monitoring is obviously constrained by the amount of available in-site sensors. Indeed, a real in-site instrumentation is largely restricted in order to remain as discrete as possible for the occupants. Thus, it is crucial to be able to propose an accurate association between efficient monitoring strategies and a low intrusive instrumentation. Then, the main work will be to assess a list of available sensors influencing the natural ventilation and absolutely required for the implementation of an efficient monitoring allowing to reach the thermal comfort.

The first step consists in optimizing numerically the building envelope in order to reach both a visual and a thermal comfort. For this first case, the natural ventilation contribution is not taken into account. The work focuses on the windows characteristics with notably the frame and glazing thermal coefficient and the solar factor. Thus, two aspects are taken into account with the thermal and the visual comfort. In order to reach those objectives, two successive sensitivity analysis methods will be employed. The qualitative Morris method is first used on an exhaustive list gathering all the input parameters to be optimized. This first method is chosen for its ability to restrain the amount of influential parameters. Output indicators for each objectives are defined with attention as scalar values in order to simplify the interpretation of the Morris method. Finally, it highlights the less influential parameters and lead to reduce the input parameters list. Thereafter, the more quantitative SRC (Standard Regression Method) is used on the restrained number of influential parameters. This method is quite restrictive given that it suggested that the building behavior can be described linearly according the variables parameters. Nevertheless, in a first approximation, many authors showed that this assumption is validated and especially in this specific case for which the natural ventilation is not considered. Finally, thanks to these sensitivity analysis methods combined with a building expertise, optimal characteristics for each window of the demonstration building are determined.

This second step gives the means to provide effective monitoring strategies of the natural ventilation in our demonstration building. More specifically, this study consists through numerical simulations to understand the influence of each measured parameters already in place on the natural ventilation and to deduce a restricted list of available sensors that would take part in the monitoring. Contrary to the previous linear case, the natural ventilation is a complex nonlinear phenomenon. As we mentioned previously, it is an aspect that requests a large attention for the choice of the sensitivity analysis methods. The Morris method allows precisely to isolate the linear from the nonlinear parameters. This technique is used on all already available measured parameters to keep only the most influential. For a better validation, the Morris method is associated with another more quantitative method. The SRC method gets less used to this application given the high non linearity of the natural ventilation. Nevertheless, many other sensitivity analysis methods could be used as the chaos polynomial or else the FAST method. Finally, a restricted list of the available sensors contained in the instrumentation already in place is retained for implementing efficient monitoring strategies of the natural ventilation in our demonstration building.

To conclude, this study applies sensitivity analysis methods to building applications. More precisely, the objective is to determine efficient passive cooling solutions allowing to reduce the solar gains and to evacuate the internal loads in an office building located in the urban tropical environment of the Reunion Island. In order to do that, two applications are specifically considered: the multi-criterion optimization of the building envelope allowing to minimize the solar gains with maintaining a suitable visual comfort and the assessment of a restrictive list of available in-site sensors that would take part in the monitoring strategies of the natural ventilation. Given the different degree of complexity, sensitivity analysis methods have to be well-chosen in order to be able to solve efficiently and accurately each problematics. In perspective, the sensitivity analysis methods could be coupled with uncertainties analysis in order to improve the robustness of the monitoring strategy. Moreover, our methodology could be completed by an additional step consisting to propose the addition of new required sensors serving the improvement of the monitoring strategies.

EnergyPlus Laboratory for Sensitivity and Uncertainty Analysis in Building Energy Modeling : The EPLab Software

Mickael Rabouille¹, Jeanne Goffart¹

¹ PUCPR, Curitiba, Brésil

Abstract

The sensitivity and uncertainty analysis in building energy modeling is a relevant topic for improvement of the buildings, the simulation models and their results.

The building physics modeling must be a tradeoff between the accuracy and a growing number of uncertain inputs. Moreover, it is a multiphysics simulation and a wide range of aspects and kind of results can be obtained in the same simulation. There is a need of reliability and transparence of the building performance simulation.

The software EPLab (EnergyPlus Laboratory for Sensitivity and Uncertainty Analysis in Building Energy Modeling) has been develop to apply uncertainty and sensitivity analysis in building performance simulation and allows wide possibilities of adaptability according to the definition of the study. The EPLab Software is coded in Matlab, it consists of the generation of sample, then the creation of idf file in order to propagate the sample by the automated call of EnergyPlus and finally the extraction and shaping of the results from the EnergyPlus outputs. The Figure 1 presents the architecture of the EPLab software.

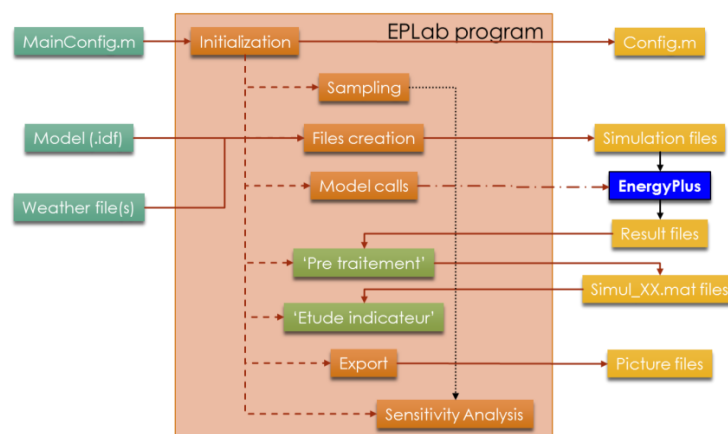


Figure 1: EPLab sotware architecture

The inputs can be fixed or statistically defined as Uniform, Gaussian with mean and variability or Discrete uniform. To evaluate the different inputs combinations, various sampling strategy can be used, the user can choose among the simple Monte-Carlo, the more relevant LHS with minimax or minimean strategies or the semi-random sequences like Halton and LP- π .

Then according to the base EnergyPlus model definition file, EPLab produces a new building model version (.idf file) for each combination of the input values from the sample. This base file is a fully functional EnergyPlus model with additional comment lines the user can produce a complex modification according to the specific inputs values for this simulation. Then the EP engine is called with the desired number of CPU core for parallel computing with one or several computers.

A verification of the completion of each simulation is made and the results are extracted and saved in a MatLab formatted file. A shaping procedure analyses in different ways the building behavior (time range, surface type, orientation, boundary condition, ...) and produces scalar or time dependent results. Finally, the results can be saved and analyzed with graphical or sensitivity methods like RBD-FAST or other sensitivity methods implemented.

The aim of the contribution is to present the ability of the EPLab software dedicated to one of the most used building performance simulation software: Energyplus. EPLab is open source and available in Github repository (Rabouille et al., 2015). Two applications of the software will be presented. The First application is about an uncertainty and sensitivity analysis applied to hygrothermal simulation of a brick building in the hot and humid climate of Singapore (J. Goffart et al. 2015). The outputs of interest are the cooling energy demand and the temporal profile of indoor humidity and temperature, the results are obtained with the four levels of wall transfer model and two assumptions for the most complex model HAMT (see Figure 2). Then another application is exposed about the uncertainty and sensitivity analysis of the calculation of the solar incident radiation on the building exterior surfaces (A.P. Almeida Rocha et al. 2016). The output of interest is the solar fraction of the south facade (see Figure 3).

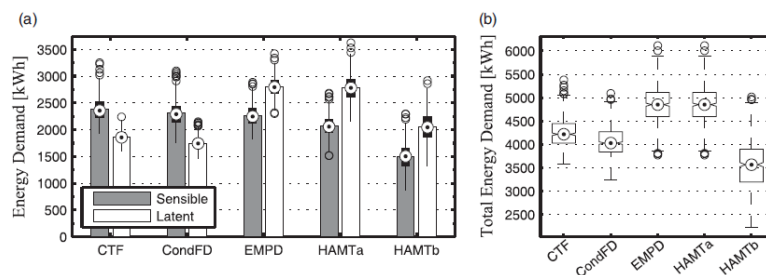


Figure 2: Box plot of the dispersion of the 600 simulations for the study on the annual results of cooling energy demand for the five cases: (a)sensible and latent cooling loads and (b) total cooling energy demand.

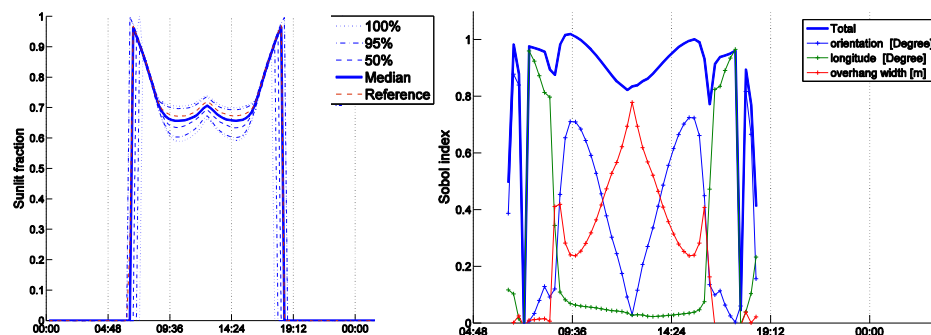


Figure 3: Solar fraction and distribution of uncertainties for the 21 march in New York at left and at right sensitivity index of the temporal profile associated with the solar fraction uncertainties (A.P. Almeida Rocha et al. 2016).

Currently an automated call of DAYSIM program is under progress, in order to perform uncertainty and sensitivity analysis combined with EnergyPlus and DAYSIM for both common natural lighting and energy performance issues.

References:

- M. Rabouille, J. Goffart, EPLab: Sensitivity and Uncertainty Analysis in Building Energy Modeling, GitHub repository, V1.8 (2015). <https://github.com/mrabouille/EPLab/releases>.
- J.Goffart, M. Rabouille, N. Mendes. « Uncertainty and sensitivity analysis applied to hygrothermal simulation of a brick building in a hot and humid climate », Journal of Building Performance Simulation, December 2015, <http://dx.doi.org/10.1080/19401493.2015.1112430>
- A.P. Almeida Rocha, J. Goffart, N.Mendes « Uncertainty assessment of the calculation of the solar incident radiation on the building exterior surfaces» Energy and Buildings , under review

Dynamic sensitivity analysis method based on Gramian

S. HAMZA

Université de Lorraine, CRAN, Nancy, France

F. ANSTETT-COLLIN

Université de Lorraine, CRAN, Nancy, France

A. BIROUCHE

Université de Haute-Alsace, MIPS, Mulhouse, France

Context

A dynamic vehicle depends on various subsystems which characterize the vehicle behavior. Each subsystem is described by a mathematical model depending on a significant number of parameters. These parameters are very often uncertain due to a lack of measurements, knowledge due to expert judgment. The uncertainty in the parameters propagates through the model and manifests itself at the model output. In order to understand the vehicle behavior, it is essential to know the parameters responsible for the model output variation. Uncertainty and sensitivity analysis can help to evaluate the impact of this lack of knowledge on the model response ([1,2]). In the literature, sensitivity analysis for dynamical models is not straightforward. In this context, the work presented in this paper investigates a novel technique of global sensitivity analysis for dynamical models. The originality of the method is to use control theory tools for sensitivity analysis purposes.

Methodology

Consider a dynamical linear system presented in state space form given by:

$$\sum_{SYS} : \begin{cases} \dot{x}(t) = A(\theta)x(t) + B(\theta)u(t) \\ y(t) = C(\theta)x(t) + D(\theta)u(t) \end{cases} \quad (1)$$

where $x(t) \in \mathbb{R}^{n_x}$ is the state vector, $y(t) \in \mathbb{R}^{n_p}$ the output vector, $u(t) \in \mathbb{R}^{n_m}$ the input vector and $t \in \mathbb{R}^+$ refers to the time. $A(\cdot) \in \mathbb{R}^{n_x \times n_x}$ is the state matrix, $B(\cdot) \in \mathbb{R}^{n_x \times n_p}$ the input matrix, $C(\cdot) \in \mathbb{R}^{l \times n_x}$ the output matrix and $D(\cdot) \in \mathbb{R}^{l \times n_p}$ the feedforward matrix. The vector $\theta = [\theta_1, \dots, \theta_{n_\theta}]$ represents the n_θ uncertain parameters. As the parameters θ_i are uncertain, they are considered as random variables defined by their probability density function (uniform, Gaussian, etc.). The uncertainty of the parameters is propagated through the model on the output $y(t)$ which becomes also uncertain. The aim is to determine the most influential parameters θ_i on the output uncertainty.

The proposed method is based on the analysis of the system energy required to drive a state to a final one by the input. If this energy is minimal, the system is said controllable. Controllability means that the system dynamics can be modified when acting on the input signal $u(t)$. The system energy depends on the uncertain parameters θ . Intuitively, if the parameter variation leads to a significant variation of the energy, it means that this parameter variation acts on the system dynamics and leads to a system that is more or less controllable. In this case, the system controllability is sensitive to this parameter variation and thus this parameter is influential on the system states. The system energy can be determined through the reachability Gramian ([3]).

The infinite time reachability Gramian, denoted for short $W_R(\theta)$, when $t_f \rightarrow \infty$ is given by:

$$W_R(\theta) = \lim_{t \rightarrow \infty} W_R(\theta, t_0, t_f) = \int_0^\infty e^{A(\theta)t} B(\theta) B(\theta)^T e^{A(\theta)^T t} dt \quad (2)$$

and it is obtained by solving the continuous-time Lyapunov equation :

$$A(\theta)W_R(\theta) + W_R(\theta)A(\theta)^T + B(\theta)B(\theta)^T = 0 \quad (3)$$

The minimal energy allowing to bring a system to a final state x_f , since $x(t_0) = 0$, is given by:

$$\|u\|_2 = \int_0^\infty u^T(t)u(t)dt = x_f^T W_R(\theta)^{-1} x_f \quad (4)$$

The quantity $x_f^T W_R^{-1} x_f$ in (4) represents an hyperellipsoid which includes all the reachable states obtained from the optimal input sequence u_{opt} . This quantity depends on the inverse of the reachability Gramian $W_R^{-1}(\theta)$. Each eigenvalue of $W_R^{-1}(\theta)$, denoted λ_i , corresponds to one system state. In fact, these eigenvalues determine the size of the axes of the hyperellipsoid and the eigenvectors determine its directions. Intuitively, the variation of an influential parameter will lead to a significant change of the dimension of the hyperellipsoid axes (see FIGURE 1).

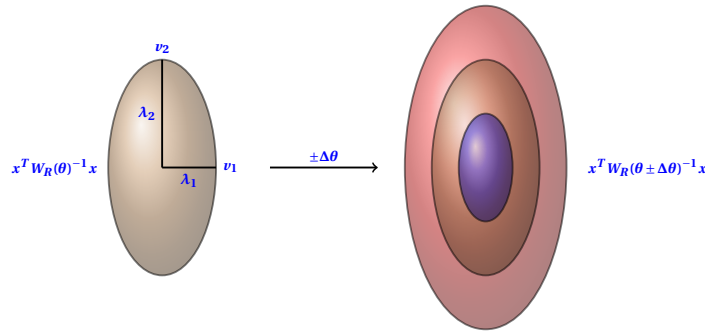


FIGURE 1 – Example of second-order system energy variation according to parameters variation.

According to (2) and (4), the eigenvalues of W_R^{-1} provide information on how controllable the system is. Higher the eigenvalues are, lower the required energy is and thus more controllable the system is ([3]). If any $\lambda_i = 0$, the system is not controllable. In fact, each eigenvalue λ_i is a function of the system parameters θ . The eigenvalues represent a measure of the controllability, that is the sensitivity of the system dynamics to variation. In this way, the parameters involved in the expression of the eigenvalues are influential on the system states. If the eigenvalue does not depend on a given parameter, this parameter is not influential on the state variation. From the structural expression of $W_R^{-1}(\theta)$, the qualitative influence of the parameters can be deduced. Then, to quantify the individual contribution of each parameter to the variance of the energy $\|u\|_2$, the Sobol' indices can be computed.

The advantage of the proposed approach is that the system energy does not depend on time and is thus scalar. Furthermore, it takes into account the global model behavior.

The approach is applied on a bicycle model describing dynamic behavior of the vehicle.

References

- [1] A. Saltelli and K. Chan and E. M. Scott, "Sensitivity analysis", Wiley, 2000.
- [2] E. De Rocquigny and N. Devictor and S. Tarantola, "Uncertainty in industrial practice", Wiley, 2008.
- [3] B. Moore, Principal component analysis in linear systems : Controllability, observability, and model reduction. Automatic Control, IEEE Transactions on (1981) 17–32.

Sobol' sensitivity analysis for stochastic numerical codes

BERTRAND, IOOSS
EDF Lab Chatou, France

THIERRY, KLEIN
Institut de Mathématiques de Toulouse, France

AGNÈS, LAGNOUX
Institut de Mathématiques de Toulouse, France

Most of the results in sensitivity analysis consider deterministic computer codes, that is codes providing the same output values for the same input variables (Iooss and Lemaître, 2015). For instance, the sensitivity indices of Sobol makes it possible to know the part of the variance output explained by each of the model input. Formally, let us consider the model

$$Y_{\text{det}} = g(X),$$

where $X = (X_1, \dots, X_p)$ is a random vector of independent input parameters (for $j = 1, \dots, p$, X_j belongs to some probability space $\mathcal{X}_j$), $Y_{\text{det}} \in \mathbb{R}$ is the code output and $g(\cdot)$ is a deterministic function representing the computer code.

In this work, we propose to deal with a stochastic computer code denoted by

$$Y_{\text{stoch}} = f(X, \varepsilon),$$

where $f(\cdot)$ is the computer code and ε is a random variable representing the physical system randomness (see Marrel et al., 2012, for a typological description of this kind of models). When performing a Sobol' sensitivity analysis on such a code, two different situations occur:

1. We are interested by the full probability density function (pdf) of the outputs. Transformation of this pdf to a few scalar quantities of interest (*e.g.* the first statistical moments of the studied variable) is a first simple solution, while metrics between pdfs can also be used (Douard and Iooss, 2013). Aggregated Sobol' indices (Gamboa et al., 2013) propose a more elegant solution as shown in Le Gratiet et al. (2016) on an application involving probability of detection curves (which look like cumulative distribution functions).
2. We are only interested by the mean value relative to the inherent randomness of the code. In this case (called "Monte Carlo calculation codes" in several engineering domains), we substitute the code by its empirical mean (called "simulator" in this paper).

We focus our analysis on the second situation. In this context, the computer code does not provide the true value of the model (denoted by $g(\cdot)$) at x but instead a value $f(x, \varepsilon)$ where ε represents the physical system randomness. A standard technique assumes that ε is a random variable such that $\mathbb{E}(f(x, \varepsilon)^2) < \infty$. Hence the real value of g at x can be represented as

$$Y = g(X) := g(X_1, \dots, X_p) = \mathbb{E}(f(X, \varepsilon)|X). \quad (1)$$

For deterministic computer code, by assuming that Y is square integrable and $\text{Var}Y \neq 0$, the corresponding vector of closed Sobol' indices is then

$$S_{\text{Cl}}^{\mathbf{u}}(g) := \left(\frac{\text{Var}(\mathbb{E}(Y|X_j, j \in u_1))}{\text{Var}(Y)}, \dots, \frac{\text{Var}(\mathbb{E}(Y|X_j, j \in u_k))}{\text{Var}(Y)} \right), \quad (2)$$

where $\mathbf{u} := (u_1, \dots, u_k)$ are k subsets of $I_p := \{1, \dots, p\}$. For X and for any subset v of I_p we denote by X^v the vector such that $X_j^v = X_j$ if $j \in v$ and $X_j^v = X'_j$ if $j \notin v$ where X' and X are two independent and identically distributed vectors. Taking two independent samples $(X_{(i)})_{i=1 \dots N}$

and $(X'_{(i)})_{i=1\dots N}$, where N is the elementary samples size, the Janon-Monod estimator of $S_{\text{Cl}}^{\mathbf{u}}(g)$ is then defined as (Janon et al., 2014):

$$T_{N,\text{Cl}}^{\mathbf{u}}(g) = \left(\frac{\frac{1}{N} \sum Y_{(i)} Y_{(i)}^{u_1} - \left(\frac{1}{2N} \sum (Y_{(i)} + Y_{(i)}^{u_1}) \right)^2}{\frac{1}{N} \sum M_{(i)}^{\mathbf{u}} - \left(\frac{1}{N} \sum Z_{(i)}^{\mathbf{u}} \right)^2}, \dots, \frac{\frac{1}{N} \sum Y_{(i)} Y_{(i)}^{u_k} - \left(\frac{1}{2N} \sum (Y_{(i)} + Y_{(i)}^{u_k}) \right)^2}{\frac{1}{N} \sum M_{(i)}^{\mathbf{u}} - \left(\frac{1}{N} \sum Z_{(i)}^{\mathbf{u}} \right)^2} \right), \quad (3)$$

$$\text{with } Y^v := g(X^v), \quad Z_{(i)}^{\mathbf{u}} = \frac{1}{k+1} \left(Y_{(i)} + \sum_{j=1}^k Y_{(i)}^{u_j} \right), \quad M_{(i)}^{\mathbf{u}} = \frac{1}{k+1} \left(Y_{(i)}^2 + \sum_{j=1}^k (Y_{(i)}^{u_j})^2 \right).$$

As the computer code cannot provide values of g , we use the Sobol' indices associated to f (instead of g) and study whether they are close to the Sobol' indices of g or not. It is then natural to replicate the code and approximate $\mathbb{E}(f(X, \varepsilon) | X = x)$ by its empirical mean. Thus we consider replicating the code output n times and we define what we call a simulator

$$\tilde{Y} := \tilde{g}(X, \varepsilon, n) = \frac{1}{n} \sum_{i=1}^n f(x, \varepsilon_{(i)}) = g(X) + \delta_n(X, \varepsilon),$$

where n is the replication number and $\delta_n(x, \varepsilon)$ is the perturbation. We then define the Sobol' indices associated to $\tilde{g}$ and their estimators using Eqs. (2) and (3) with $\tilde{Y}$ instead of Y . Moreover, we prove that the estimator $T_{N,\text{Cl}}^{\mathbf{u}}(\tilde{g})$ can be used to approximate the true Sobol' indices $S_{\text{Cl}}^{\mathbf{u}}(g)$. Indeed, following the proofs of Janon et al. (2014), we derive a Central Limit Theorem for this estimator (not shown here in this short abstract).

We also numerically study the convergence of the Sobol' indices estimates with respect to the replication number n and sample size N (size of the elementary samples for Sobol' estimates) considering the following toy function:

$$f(X_1, X_2, \varepsilon) = \sin(X_1(\varepsilon_1 + \varepsilon_2 X_2)) + \varepsilon_3,$$

with the independent random variables $X_1 \sim \mathcal{U}[0, 1]$, $X_2 \sim \mathcal{U}[0, 1]$, $\varepsilon_1 \sim \mathcal{N}(1, 1)$, $\varepsilon_2 \sim \mathcal{N}(2, 1)$ and $\varepsilon_3 \sim \mathcal{U}[0, 1]$. This leads to a function g defined by

$$g((x_1, x_2)) = \mathbb{E}(f(X_1, X_2, \varepsilon) | X_1 = x_1, X_2 = x_2) = \frac{1}{2} + \sin(x_1(1 + 2x_2))e^{-\frac{x_1^2}{2}(1+x_2^2)}.$$

Finally, the results will be applied to a Monte Carlo simulator of industrial asset management strategies, where the variable of interest is an economic indicator (the Net Present Value).

References:

- F. Douard and B. Iooss (2013), Dealing with uncertainty in technical and economic studies of investment strategy optimization. *Proceedings of MASCOT-SAMO 2013 Conference*, Nice, France, July 2013.
- F. Gamboa, A. Janon, T. Klein and A. Lagnoux (2014), Sensitivity analysis for multidimensional and functional outputs. *Electronic Journal of Statistics*, 8:575-603.
- B. Iooss and P. Lemaître (2015), A review on global sensitivity analysis methods. In *Uncertainty management in Simulation-Optimization of Complex Systems: Algorithms and Applications*, C. Meloni and G. Dellino (eds), Springer.
- A. Janon, T. Klein, A. Lagnoux, M. Nodet and C. Prieur (2014), Asymptotic normality and efficiency of two Sobol index estimators. *ESAIM: Probability and Statistics*, 18:342-364.
- L. Le Gratiet, B. Iooss, T. Browne, G. Blatman, S. Cordeiro and B. Goursaud (2016). Model assisted probability of detection curves: New statistical tools and progressive methodology, *Submitted*.
- A. Marrel, B. Iooss, S. da Veiga and M. Ribatet (2012), Global sensitivity analysis of stochastic computer models with joint metamodels. *Statistics and Computing*, 22:833-847.

[Bertrand Iooss; EDF R&D, 6 Quai Watier, 78401 Chatou, France]

[bertrand.iooss@edf.fr – <http://www.gdr-mascotnum.fr/doku.php?id=iooss1>]

Numerical stability of Sobol' indices estimation formula

MICHAEL, BAUDIN
EDF Lab Chatou, France

KHALID, BOUMHAOUT
EDF Lab Chatou, France

THIBAUT, DELAGE
EDF Lab Chatou, France

BERTRAND, IOOSS
EDF Lab Chatou, France

JEAN-MARC, MARTINEZ
CEA Saclay, France

Variance-based sensitivity analysis has become a common practice when using computer models in engineering studies (Ferretti et al., 2016). The so-called Sobol' indices express the share of the model output variance that is due to a given model input or input combination and write for instance (Sobol, 1993; Saltelli, 2002)

$$S_i = \frac{V_i}{V} = \frac{\text{Var}[\mathbb{E}(G(X)|X_i)]}{\text{Var}[G(X)]} \text{ and } S_i^{\text{tot}} = \frac{V_i^{\text{tot}}}{V} = 1 - \frac{V_{-i}}{V} = 1 - \frac{\text{Var}[\mathbb{E}(G(X)|X_{-i})]}{\text{Var}[G(X)]}, \quad (1)$$

where $G(X)$ is the computer model, $X = (X_1, \dots, X_d) \in \mathbb{R}^d$ are the model inputs (independent random variables), $i = 1, \dots, d$, and X_{-i} is the input vector except X_i . S_i , the first-order Sobol' index, only includes the sole effect of X_i , while S_i^{tot} , the total Sobol' index, takes into account all the effects of X_i including its interaction effects with other inputs. For a direct estimation (without additional modeling), several sampling-based formulas have been proposed in the literature (see Prieur and Tarantola, 2017, for a recent review), but have shown some instabilities from a numerical point of view, delivering different behavior in different cases.

We focus on the estimators which provide $(\hat{S}_i, \hat{S}_i^{\text{tot}})$, estimates of (S_i, S_i^{tot}) , by using two independent input designs $\mathbf{A}$ and $\mathbf{B}$, matrices with n rows (sample size) and d columns:

- Sobol-Saltelli estimator (Sobol, 1993; Saltelli, 2002):

$$\hat{V}_i = \frac{1}{n-1} \sum_{k=1}^n G(\mathbf{B}^{(k)})G(\mathbf{A}_{B(i)}^{(k)}) - \hat{G}_1^2; \quad \hat{V}_{-i} = \frac{1}{n-1} \sum_{k=1}^n G(\mathbf{A}^{(k)})G(\mathbf{A}_{B(i)}^{(k)}) - \hat{G}_0^2, \quad (2)$$

where $\mathbf{A}_{B(i)}$ is a re-sampled matrix, where all columns come from $\mathbf{A}$ except column i which comes from $\mathbf{B}$, and where the two square means and the variance are estimated by

$$\hat{G}_0^2 = \left[\frac{1}{n} \sum_{k=1}^n G(\mathbf{A}^{(k)}) \right]^2, \quad \hat{G}_1^2 = \frac{1}{n} \sum_{k=1}^n G(\mathbf{A}^{(k)})G(\mathbf{B}^{(k)}), \quad \hat{V} = \frac{1}{n-1} \sum_{k=1}^n G(\mathbf{A}^{(k)})^2 - \frac{n}{n-1} \hat{G}_0^2. \quad (3)$$

- Mauntz estimator (Mauntz, 2002):

$$\hat{V}_i = \frac{1}{n-1} \sum_{k=1}^n G(\mathbf{B}^{(k)}) \left(G(\mathbf{A}_{B(i)}^{(k)}) - G(\mathbf{A}^{(k)}) \right); \quad \hat{V}_i^{\text{tot}} = \frac{1}{n-1} \sum_{k=1}^n G(\mathbf{A}^{(k)}) \left(G(\mathbf{A}^{(k)}) - G(\mathbf{A}_{B(i)}^{(k)}) \right). \quad (4)$$

- Jansen estimator (Jansen, 1999):

$$\hat{V}_i = \hat{V} - \frac{1}{2n-1} \sum_{k=1}^n \left(G(\mathbf{B}^{(k)}) - G(\mathbf{A}_{B(i)}^{(k)}) \right)^2; \quad \hat{V}_i^{\text{tot}} = \frac{1}{2n-1} \sum_{k=1}^n \left(G(\mathbf{A}^{(k)}) - G(\mathbf{A}_{B(i)}^{(k)}) \right)^2. \quad (5)$$

- Martinez estimator (Martinez, 2011): By noticing that

$$S_i = \rho(G(\mathbf{B}), G(\mathbf{A}_{B(i)})) \text{ and } S_i^{\text{tot}} = 1 - \rho(G(\mathbf{A}), G(\mathbf{A}_{B(i)})) \quad (6)$$

where ρ is the linear correlation coefficient, the Sobol' indices can be estimated using the well-conditioned empirical formula of ρ (*i.e.* using the product of differences).

Remark 1: The denominator of the indices, the model variance V , can be estimated from several ways. We restrict our study to the one of Eq. (3).

Remark 2: The same $n(d+2)$ evaluations are needed for applying the four estimators (2), (4), (5) and (6). A direct comparison between them, using a bootstrap technique to obtain confidence intervals, is possible in practice if $\mathbf{A}$ and $\mathbf{B}$ are i.i.d samples.

Remark 3: For the Martinez estimator, asymptotic confidence intervals are approximated via a Fisher's transformation applied to the sample correlation coefficients $\hat{S}_i$ and $\hat{S}_i^{\text{tot}}$ from Eq. (6). For the classical 95% confidence level, we have:

$$\text{Prob}(S_i \in [\tanh(\frac{1}{2} \ln \frac{1 + \hat{S}_i}{1 - \hat{S}_i} - \frac{1.96}{\sqrt{n-3}}), \tanh(\frac{1}{2} \ln \frac{1 + \hat{S}_i}{1 - \hat{S}_i} + \frac{1.96}{\sqrt{n-3}})]) \simeq 0.95, \quad (7)$$

$$\text{Prob}(S_i^{\text{tot}} \in [1 - \tanh(\frac{1}{2} \ln \frac{1 + \hat{S}_i^{\text{tot}}}{1 - \hat{S}_i^{\text{tot}}} + \frac{1.96}{\sqrt{n-3}}), 1 - \tanh(\frac{1}{2} \ln \frac{1 + \hat{S}_i^{\text{tot}}}{1 - \hat{S}_i^{\text{tot}}} - \frac{1.96}{\sqrt{n-3}})]) \simeq 0.95. \quad (8)$$

It is only valid under Gaussian hypothesis of the output variable distribution. Current works aim at extending this result to non-Gaussian distribution (Touati, 2016).

In this communication, two pathological issues of the estimators' behavior are studied:

1. Non-centered output. In this case, we show that the Sobol-Saltelli and Mauntz estimators are subject to a non-negligible bias, while the other estimators are insensitive to this effect.
2. Small sensitivity indices. In this case, the numerical precision obtained for the Sobol' indices depend on the conditioning of each estimator formula. Indeed, when the terms are close to zero, differences between products (as in the Sobol-Saltelli estimator) are more sensitive than products of differences.

Numerical studies will illustrate all these effects for the different estimators, demonstrating that the Martinez estimator is particularly robust.

References:

- F. Ferretti, A. Saltelli and S. Tarantola (2016), Trends in sensitivity analysis practice in the last decade, *Science of the Total Environment*, 568:666-670.
- M.J.W. Jansen (1999), Analysis of variance designs for model outputs, *Computer Physics Communication* 117:35-43.
- J-M. Martinez (2011), Analyse de sensibilité globale par décomposition de la variance, *Presentation in "Journée des GdR Ondes & Mascot Num"*, 13 janvier 2011, Institut Henri Poincaré, Paris, France.
- W. Mauntz (2002). *Global sensitivity analysis of general nonlinear systems*. Master's thesis, Imperial College.
- C. Prieur and S. Tarantola (2017). Variance-based sensitivity analysis: Theory and estimation algorithms. In: *Springer Handbook on UQ*, R. Ghanem, D. Higdon and H. Owadi (Eds).
- A. Saltelli (2002), Making best use of model evaluations to compute sensitivity indices, *Computer Physics Communication*, 145:580-297.
- I. Sobol (1993), Sensitivity estimates for non linear mathematical models. *Mathematical Modelling and Computational Experiments*, 1:407-414.
- T. Touati (2016), Confidence intervals for Sobol' indices. *Proceedings of the SAMO 2016 Conference*, Réunion Island.

[Bertrand Iooss; EDF R&D, 6 Quai Watier, 78401 Chatou, France]

[bertrand.iooss@edf.fr – <http://www.gdr-mascotnum.fr/doku.php?id=iooss1>]

SOBOL' INDICES FOR PROBLEMS DEFINED IN NON-RECTANGULAR DOMAINS

S. Kucherenko, O.V. Klymenko, N. Shah
Imperial College London, London, SW7 2AZ, UK
s.kucherenko@imperial.ac.uk

Uncertainty and sensitivity analysis has been recognized as an essential part of model applications. Global sensitivity analysis (GSA) is used to identify key parameters whose uncertainty most affects the output. This information can be used to rank variables, fix or eliminate unessential variables and thus decrease problem dimensionality. Among different approaches to GSA variance-based Sobol sensitivity indices (SI) are most frequently used in practice owing to their efficiency and ease of interpretation [1-3]. Most existing techniques for GSA were designed under the hypothesis that model inputs are independent. However, in many cases there are dependences among inputs, which may have significant impact on the results. Such dependences in a form of correlations have been considered in the generalised Sobol GSA framework developed by Kucherenko *et al*, [4]. However, there is an even wider class of models involving inequality constraints (which naturally leads to the term constrained GSA or cGSA) imposing structural dependences between model variables. This implies that the parameter space may no longer be considered to be an n -dimensional hypercube which is the case in existing GSA methods, but may assume any shape depending on the number and nature of constraints. This class of problems encompasses a wide range of situations encountered in the natural sciences, engineering, design, economics and finances where model variables are subject to certain limitations imposed e.g. by conservation laws, geometry, costs, quality constraints etc.

The development of efficient computational methods for cGSA is challenging because of potentially arbitrary shape of the feasible domain of model variables variation, thus requiring the development of special Monte Carlo or quasi-Monte Carlo sampling techniques and methods for computing sensitivity indices. We developed a novel method for estimation of Sobol' SI for models $f(x_1, \dots, x_n)$ defined in a non-rectangular domain Ω^n . Consider an arbitrary subset of the variables $y = (x_{i_1}, \dots, x_{i_s})$, $1 \leq s < n$ and a complementary subset $z = (x_{i_{s+1}}, \dots, x_{i_n})$, so that $(x_1, \dots, x_n) = (y, z)$. Then formulas for the main effect and total Sobol' SI have the following form:

$$S_y = \frac{1}{D} \left[\int_{\Omega^n} f(y', z') p^\Omega(y', z') dy' dz' \left[\int_{\Omega^{n-s}} \frac{f(y', z)}{p^\Omega(y')} p^\Omega(y', z) dz - \int_{\Omega^n} f(y, z) p^\Omega(y, z) dy dz \right] \right],$$

$$S_y^T = \frac{1}{2D} \int_{\Omega^n} \int_{\Omega^s} [f(y, z) - f(y', z)]^2 p^\Omega(y, z) \frac{p^\Omega(y', z)}{p^\Omega(z)} dy dy' dz.$$

Here $p^\Omega(y, z)$ is a joint probability distribution and $p^\Omega(y)$ is a marginal distribution. Both distributions are defined in Ω^n . We propose two methods for estimation Sobol' SI: 1) quadrature integration method which may be very efficient for problems of low and medium dimensionality; 2) MC/QMC estimators based on the acceptance-rejection sampling method. A few model test functions with constraints are considered for which we found analytical solutions. These solutions are used as benchmark test for verifying for the quadrature and MC and QMC integrations methods. One of the models is the K-function

$K = \sum_{i=1}^n (-1)^i \prod_{j=1}^i x_j$, where variables x_j , $j=1, \dots, n$, $n=4$ are independent uniformly distributed random variables in $[0, 1]$. We consider four different cases for domain definitions. The first one is an unconstrained problem ($x \in H^n$). In the other three cases the unit hypercube is divided by a hyperplane into two parts one of which is the permissible region for the problem variables x_j , $j=1, \dots, n$. The constraints are as follows:

$$I_1 : x_1 + x_2 \leq 1,$$

$$I_2 : x_3 + x_4 \leq 1,$$

$$I_3 : x_1 + x_3 \leq 1.$$

These constraints can be represented using the following indicator functions: $I_1 = U(1 - x_1 - x_2)$, $I_2 = U(1 - x_3 - x_4)$, $I_3 = U(1 - x_1 - x_3)$. Fig. 1 shows a schematic plot illustrating I_1 constraint in the 3D space.

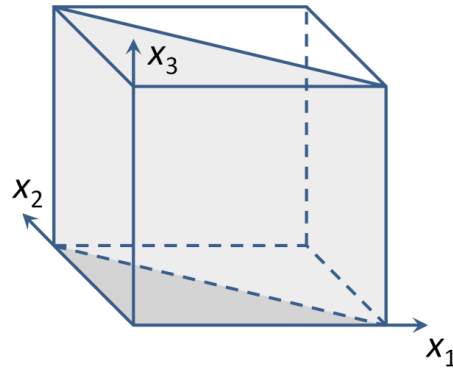


Fig. 1. Schematic representation of a permissible region for the K-function (shaded area) in the 3D case.

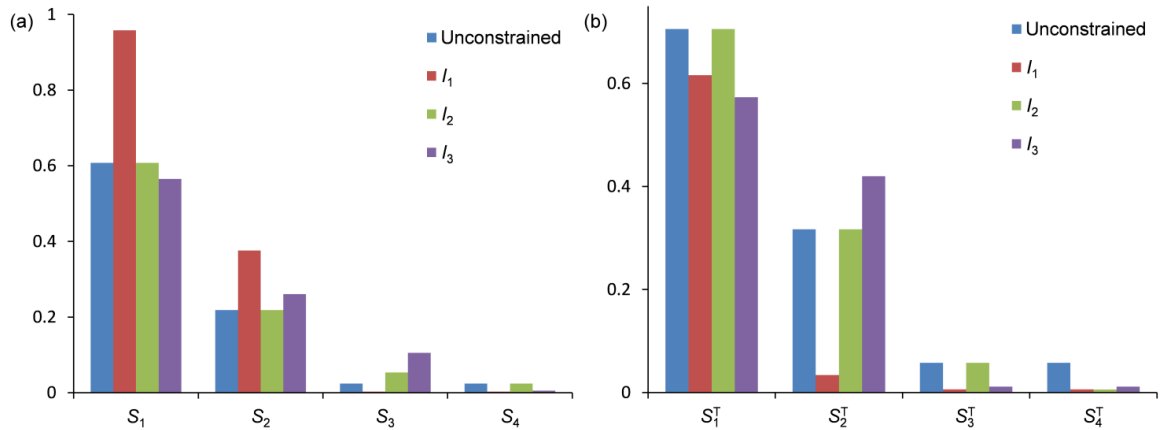


Fig. 2. (a) Main effect and (b) total sensitivity indices of the K-function in H^4 for the unconstrained and constraints cases

The values of S_i and S_i^T for all four cases are presented in Fig. 2. The method is shown to be general and efficient.

1. A. Saltelli, S. Tarantola, F. Campolongo, M. Ratto, *Sensitivity analysis in practice*. London, Wiley; 2004.
2. I. Sobol', Global sensitivity indices for nonlinear mathematical models and their Monte Carlo estimates. *Math Comput Simul*, 55(1-3), 271-280, 2001.
3. A. Owen, Better estimation of small Sobol' sensitivity indices. *ACM Trans on Mod and Comp Simul*, 23(11), 1-17, 2013.
4. S. Kucherenko, S. Tarantola, P. Annoni. Estimation of global sensitivity indices for models with dependent variables. *Comput. Phys. Commun.* 183, 937-946, 2012.

GLOBAL SENSITIVITY ANALYSIS OF NON-DOMESTIC BUILDINGS THERMAL BEHAVIOUR

S. Kucherenko, F. Abdoussi, A. Hirvoas, B. Howard
Imperial College London, London, SW7 2AZ, UK
s.kucherenko@imperial.ac.uk

Non-domestic buildings are usually equipped with a centralised energy management system, the Building Energy Management System (BEMS). The optimisation of the control strategy and the implementation of advanced algorithms in BEMS would allow substantial energy savings in the operation of the energy systems such as Heating, Ventilation and Air Conditioning (HVAC) plants. Simulation tools are necessary to quantify the potential savings.

The existing building simulation tools do not allow the simulation of all advanced control strategies or all interactions between the indoor environment, the HVAC plant and the BEMS. An integrated model of a building was developed in MATLAB in order to perform control strategy assessments. This model comprises three submodels: a model of the building thermal behaviour, a model of the HVAC plant and the control algorithms of the BEMS. In order to be suitable for control strategy evaluation, the model should comply with the following specifications: it should be dynamic, have a short computational time and be replicable to many buildings. The modelling choices result from these specifications.

The building thermal behaviour is simulated through a thermal network model, presented here in a matrix form:

$$\begin{bmatrix} \frac{dT_i}{dt} \\ \frac{dT_w}{dt} \end{bmatrix} = \begin{bmatrix} -\left(\frac{1}{R_{air}} + \frac{1}{R_{eq}}\right) \cdot \frac{1}{C_{air}} & \frac{1}{R_{air} \cdot C_{air}} \\ \frac{1}{R_{air} \cdot C_{envelope}} & -\left(\frac{1}{R_{air}} + \frac{1}{R_{envelope}}\right) \cdot \frac{1}{C_{envelope}} \end{bmatrix} \begin{bmatrix} T_i \\ T_w \end{bmatrix} + \begin{bmatrix} \frac{1}{C_{air}} & \frac{1}{R_{eq} \cdot C_{air}} \\ 0 & \frac{1}{R_{envelope} \cdot C_{envelope}} \end{bmatrix} \begin{bmatrix} \dot{Q}_{internal} \\ T_o \end{bmatrix}$$

Here T_i and T_w are the inside and wall temperatures respectively, $\frac{1}{R_{eq}} = \frac{1}{R_{\{window, infiltration\}}} + \frac{1}{R_{ventilation}}$.

Such model has the advantage of being flexible: each thermal zone is represented by a node of the thermal network and a finer granularity is achieved by increasing the number of nodes. The structure of the thermal network is derived from heat balance equations and their physical interpretation. The values of the parameters of the thermal network model are determined by identification with measured and metered data from the modelled building. The HVAC system is modelled through a white-box model. The plant model main outputs are: the heat supplied to the indoor environment (which is an input of the building thermal submodel) and the equipment statuses (in order to have some insight about the way the HVAC system meets the heating, cooling and ventilation demands). More specifically model outputs are the evolution over time of the indoor temperature, the building thermal envelope temperature, the equipment statuses and the thermal energy provided to the indoor environment. Along with these time dependent curves, some key indicators were chosen in order to have an insight into indoor environment conditions and energy expenditure. There is a choice of the control strategies.

Uncertainty and sensitivity analysis is an important step in accessing building model applications as practically all parameters are either known with some tolerance or are uncertain. We apply variance based global sensitivity analysis (GSA) to identify key parameters whose uncertainty most affects the output. We also use this information to analyse identifiability of parameters during the calibration process.

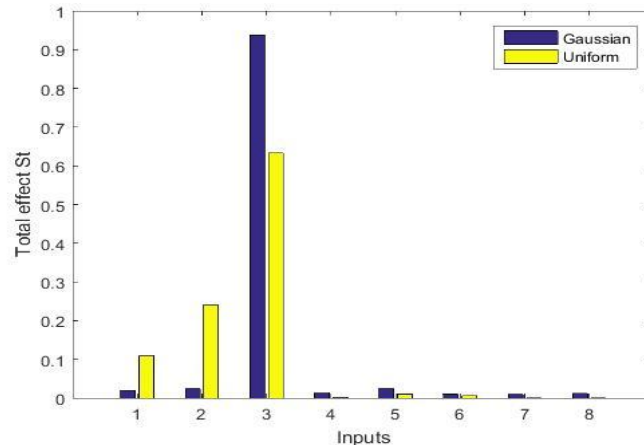
The values of conductance and resistance parameters with their uncertainty distribution parameters for a typical supermarket building in the UK are given in this table:

		Uniform	Gaussian
1	R_{air}	$R_{air_low} = 3.71e-05$ $R_{air_sup} = 4.54e-05$	mean=4.127e-5 std=2.17e-6
2	$R_{wall} = R_{envelope}$	$R_{wall_low} = 2.11e-04$ $R_{wall_sup} = 2.57e-04$	mean=2.34e-4 std=2.5e-5
3	$R_{eq} = R_{ventilation_window}$	$R_{eq_low} = 2.48e-04$ $R_{eq_sup} = 3.04e-04$	mean=2.76e-4 std=7.44e-5
4	$C_{air} = C_{air}$	$C_{air_low} = 8.32 e+07$ $C_{air_sup} = 1.02e+08$	mean=9.24e+7 std=1.42e+7
5	$C_{wall} = C_{envelope}$	$C_{wall_low} = 9.95 e+08$ $C_{wall_sup} = 1.22e+09$	mean=11.1e+8 std=2.35e+9

Two different distributions are considered: Uniform and Gaussian. For Uniform distribution 10% variation is assumed for all inputs. For the Gaussian distribution coefficients of variation (the standard deviation normalized by the expected value $CV = \frac{\sigma}{\mu}$) are presented in this table:

Inputs	Coefficient of variation(%)
1	5.25
2	10.68
3	26.96
4	15.40
5	2.12

The values of total sensitivity indices for eight inputs (five building parameters discussed above and three HVAC parameters) for the thermal energy provided to the indoor environment are presented in the following Fig.:



SobolGSA which is a general purpose GUI driven global sensitivity analysis and metamodeling software was used for GSA [1]. One can see that R_{eq} is the most important factor followed by R_{wall} and R_{air} for the case of Uniform distribution while for the case of Gaussian distribution R_{eq} is the only important factor. We also developed a full scale model using EnergyPlus software and linked it with SobolGSA. The results of the MATLAB based thermal network model are compared with those of a full scale model.

1. SobolGSA software (2016). <http://www.imperial.ac.uk/process-systems-engineering/research/free-software/sobolgsa-software/>
2. EnergyPlus software (2016) <https://energyplus.net/>

A minimum variance unbiased (generalized) estimator of total sensitivity indices: an illustration to a flood risk model

MATIEYENDOU LAMBONI^{a,b}

^a *University of Guyane (UG), Department of Science and Technology (DFRST), 97346 Cayenne, French Guiana*

^b *228-UMR Espace-Dev, 97323 Cayenne Cedex, French Guiana*

I. Objective Variance-based sensitivity analysis [1-2] and multivariate sensitivity analysis [3-5] aim at apportioning the variability of the model output(s) into input factors and their interactions. Sobol's total index, which accounts for the effects of interactions, is often used for selecting the most influential parameters. In this paper, we propose a generalized and optimal estimator of the variance of the total effect (non-normalized total sensitivity index- TSI). The generalized and optimal estimator of the non-normalized TSI makes use of p -fold sets of input values to obtain the TSI estimates. When $p = 2$, we obtain the Jansen's estimator. An illustration to a flood model shows that we can improve the TSI estimations using p larger than 2.

II. Methods Let $Y = f(\mathbf{X})$ be a model output with $\mathbf{X} = (X_1, \dots, X_d)$, d independent input factors (A1). Under assumption $\mathbb{E}(f^2(\mathbf{X})) < +\infty$ (A2), we have the Hoeffding decomposition:

$$f(\mathbf{X}) = \sum_{u \subseteq \{1, 2, \dots, d\}} f_u(X_u), \quad (1)$$

where $f_\emptyset = \mathbb{E}[f(\mathbf{X})]$, $f_j(X_j) = \mathbb{E}[f(\mathbf{X})|X_j] - f_\emptyset$, and $\mathbb{E}[f_u(X_u)] = 0$.

It is shown in [2,6] that the non-normalized TSI of a set of inputs $\mathbf{X}_u = (X_j, j \in u)$, is also defined as follows:

$$D_u^{tot} = \mathbb{E}(f(\mathbf{X}) - \mathbb{E}[f(\mathbf{X})|\mathbf{X}_{\sim u}])^2. \quad (2)$$

Definition 1 Let us consider independent samples $\mathbf{X}_u^{(1)}, \dots, \mathbf{X}_u^{(p)}$ from the measure $\mu(\mathbf{X}_u)$, $\mathbf{X}_u^{(1)'}, \dots, \mathbf{X}_u^{(p)'}$ from $\mu(\mathbf{X}_{\sim u})$ and $\mathbf{X}_{\sim u} = (X_j, j \in \{1, 2, \dots, d\} \setminus u)$ from $\mu(\mathbf{X}_{\sim u})$. We define a kernel function as:

$$K(\mathbf{X}_u^{(1)}, \dots, \mathbf{X}_u^{(p)}, \mathbf{X}_{\sim u}) = \frac{p-1}{p^2} \sum_{l=1}^p \left(\sum_{j=1}^p c_j^{(l)} [f(\mathbf{X}_u^{(l)}, \mathbf{X}_{\sim u}) - f(\mathbf{X}_u^{(j)}, \mathbf{X}_{\sim u})] \right)^2, \quad (3)$$

with $c_j^{(l)} = \frac{1}{p-1}$ if $j \neq l$ and 0 otherwise (A3).

If we define $\sigma_{l,1}^2 = \text{Cov} \left[K(\mathbf{X}_u^{(1)}, \dots, \mathbf{X}_u^{(l)}, \mathbf{X}_u^{(l+1)}, \dots, \mathbf{X}_u^{(p)}, \mathbf{X}_{\sim u}), K(\mathbf{X}_u^{(1)}, \dots, \mathbf{X}_u^{(l)}, \mathbf{X}_u^{(l+1)'}, \dots, \mathbf{X}_u^{(p)'}, \mathbf{X}_{\sim u}) \right]$, then it satisfies $\sigma_{l,1}^2 = \mathbb{V} \left(\mathbb{E} \left[K(\mathbf{X}_u^{(1)}, \dots, \mathbf{X}_u^{(p)}, \mathbf{X}_{\sim u}) | \mathbf{X}_u^{(1)}, \dots, \mathbf{X}_u^{(l)}, \mathbf{X}_{\sim u} \right] \right)$ ([7]).

Theorem 1 Let $Y = f(\mathbf{X})$ and consider independent samples $(\mathbf{X}_{i,u}^{(1)}, \mathbf{X}_{i,\sim u})$, $\dots$, $(\mathbf{X}_{i,u}^{(p)}, \mathbf{X}_{i,\sim u})$ from $\mu(\mathbf{X})$ with $i = 1, 2, \dots, m$. Under assumptions A1, A3 ($\forall j, l = 1, 2, \dots, p$, $c_j^{(l)} = \frac{1}{p-1}$ if $j \neq l$ and 0 otherwise), A4 ($\mathbb{E}[f^4(\mathbf{X})] < +\infty$), and A5 ($2 \leq p$), we have:

i) the optimal, unbiased estimator of D_u^{tot} for a given p is:

$$\hat{D}_u^{tot} = \frac{p-1}{mp^2} \sum_{i=1}^m \sum_{l=1}^p \left(\sum_{j=1}^p c_j^{(l)} [f(\mathbf{X}_{i,u}^{(l)}, \mathbf{X}_{i,\sim u}) - f(\mathbf{X}_{i,u}^{(j)}, \mathbf{X}_{i,\sim u})] \right)^2; \quad (4)$$

ii) some properties of $\hat{D}_u^{tot}$ are:

$$m \mathbb{E} \left(\hat{D}_u^{tot} - D_u^{tot} \right)^2 = \sigma_{p,1}^2, \quad \sqrt{m} \left(\hat{D}_u^{tot} - D_u^{tot} \right) \xrightarrow{\mathcal{D}} \mathcal{N}(0, \sigma_{p,1}^2). \quad (5)$$

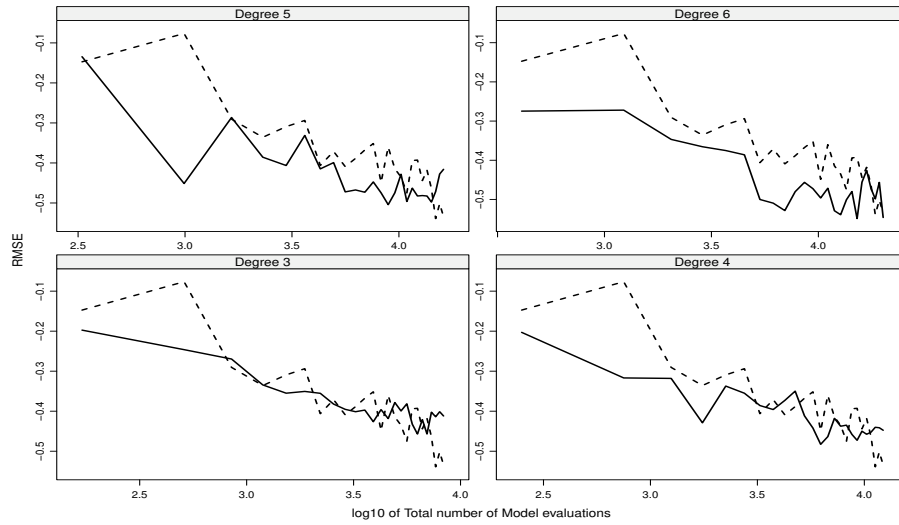


Figure 1: Log-RMSEs against the total number of model runs (in $\log_{10}$) for four values of the degree $p = 3, 4, 5, 6$. For each degree, we show the corresponding RMSE (solid line) and the RMSE for Jansens estimator (dashed line).

Proof The kernel $K(\cdot)$ is symmetric with respect to its first argument $(\mathbf{X}_u^{(1)}, \dots, \mathbf{X}_u^{(p)})$ and we have $\mathbb{E} \left[K(\mathbf{X}_u^{(1)}, \dots, \mathbf{X}_u^{(p)}, \mathbf{X}_{\sim u}) \right] = D_u^{tot}$. The ponits i) and ii) are obtained using the properties of U-statistics. (see [8] for more details). $\square$

III. Results and Conclusions To illustrate our approach, we consider a flood model that simulates the height of a river compared to the height of a dyke [2,9]. The model includes 8 input factors. We compared the TSI estimates for four different values of $p = 3, 4, 5, 6$ to those for Jansen's estimator ($p = 2$), using Sobol's design. Figure 1 shows the average of the root mean square errors (RMSEs) of the 8 inputs against the total number of model runs for each degree $p = 3, 4, 5, 6$ compared to $p = 2$. It comes out that the degree $p = 6$ provides the best estimations.

References:

- [1] Saltelli A, Ratto M, Andres T, Campolongo F, Cariboni J, Gatelli D, Salsana M, Tarantola S (2008) Global sensitivity analysis - The primer. Wiley.
- [2] Higdon, D., Ghanem, R. and Owhadi H. (2016) Handbook of Uncertainty Quantification. Springer.
- [3] Lamboni M, Makowski D, Lehuger S, Gabrielle B, Monod H (2009) Multivariate global sensitivity analysis for dynamic crop models. Fields Crop Reasearch 113: 312 - 320.
- [4] Lamboni M, Monod H, Makowski D (2011) Multivariate sensitivity analysis to measure global contribution of input factors in dynamic models. Reliability Engineering and System Safety 96: 450 - 459.
- [5] Gamboa F, Janon A, Klein T, Lagnoux A (2014) Sensitivity indices for multivariate outputs. Comptes Rendus de l'Acadmie des Sciences p In press.
- [6] Lamboni M (2013) New way of estimating total sensitivity indices. In: Proceedings of the 7th International Conference on Sensitivity Analysis of Model Output (SAMO 2013), Nice, France.
- [7] Ferguson TS (1996) A Course in Large Sample Theory. Chapman-Hall, New York.
- [8] Lamboni M (2016) Global sensitivity analysis: a generalized, unbiased and optimal estimator of total-effect variance (submitted).
- [9] Iooss B Revue sur l'analyse de sensibilité globale de modèles numériques, Journal de la Société Française de Statistique 152 (2011) 1 - 23.

[M. Lamboni; UG, DFRST, Cayenne-French Guiana, matieyendou.lamboni@gmail.com]

Sensitivity analysis and metamodeling methods for designing buffer strips to protect water from pesticide transfers.

Lauvernet Claire ⁽¹⁾, Helbert C. ⁽²⁾, Catalogne C. ⁽¹⁾, Carluer N ⁽¹⁾, Muñoz-Carpena R. ⁽³⁾

(1) *Irstea, UR MAEP, Equipe « Pollutions agricoles diffuses » - 5 rue de la Doua, CS70077 – 69626 Villeurbanne Cedex – France – claire.lauvernet@irstea.fr*

(2) *Univ Lyon, UMR CNRS 5208, École Centrale de Lyon - ICJ – 36 av. Guy de Collongue - 69134 Ecully Cedex - France*

(3) *Univ Florida, Agricultural and Biological Engineering, 287 Frazier Rogers Hall, PO Box 110570 Gainesville, FL 32611-057 - USA*

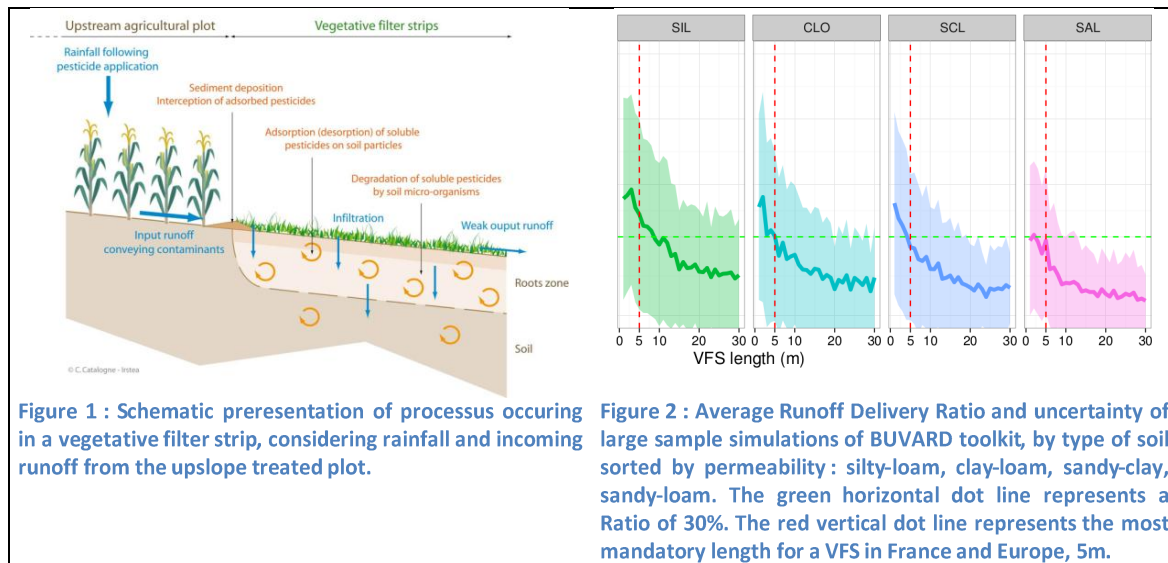
In France, significant amounts of pollutants are measured in surface water, partly due to the use of pesticides by agriculture. In order to decrease this contamination, reasoning pesticides applications at field and taking into account the rainfall forecasting are recommended in addition to more global management practices at the watershed scale. Among these, buffer zones are landscape elements that are either maintained or intentionally set up to ensure the interception and the mitigation of contaminant transfers arising from fields towards aquatic environments. Among buffer zones, Vegetative Filter Strips (VFS, such as grassed or wooded strips) can be useful tools to help achieving the "good ecological status" of the European Water Framework Directive in watersheds (see Figure 1). They are now mandatory along rivers in more and more countries, due to their recognized effectiveness to limit pesticide and sediment transfer by surface runoff (Asmussen et al., 1977; Dosskey, 2001). However, the general effectiveness of edge-of-field buffer strips to reduce runoff transport of pesticides can highly depend on many conditions of implementation and maintenance. It particularly requires an optimal design (position in the hillslope and size) which depends on the agronomic conditions, the studied hillslope's soil characteristics and the climate.

In this context, Irstea developed a methodology which allows designing site-specific VFS by simulating their efficiency to limit transfers among a hillslope (Carluer et al., 2014). The modeling toolkit BUVARD (Buffer strip for runoff Attenuation and pesticides Retention Design tool) consists in several steps like analyzing the watershed and its characteristics (soil, climate, cultural practices), and running dynamical models, in particular the mechanistic model VFSSMOD (Munoz-Carpena et al., 1999). At the end this toolkit delivers the optimal VFS width considering the needed filter efficiency (for example, 70% of runoff reduction). This very complete method assumes that the user provides detailed field knowledge and data (type of soil of the contributive area and of the VFS, rainfall rate, water table depth, slope, etc.), which are not easily available in many practical applications. Moreover, the variety of tools, which rely on several interfaces or several programming languages makes it relatively difficult to take over the design procedure.

In order to fill in the lack of real data in many practical applications, a first study was conducted on a set of virtual scenarios, among which the user would choose the most relevant one considering its own situation. They were selected by expertise to encompass a large range

of agro-pedo-climatic conditions in Europe, describing both the upslope agricultural field and the VFS characteristics. The sizing method was applied on all these scenarios, and a Global Sensitivity Analysis (GSA, Sobol method) of the VFS optimal length was performed to understand the influence of parameters and their interactions, and identify priorities for data collecting and management. The GSA applied in a specific watershed of North-West of France (Lauvernet et al 2015) showed that the optimal VFS width is very sensitive to i) the Curve Number (an empirical parameter describing the potential surface runoff generation of the contributive area, USDA-SCS, 1972), ii) the saturated hydraulic conductivity of the VFS soil and iii) the water table depth. These parameters, which particularly depend on local conditions, must then be defined as accurately as possible. The performed sensitivity analysis procedure provided interesting results, but these are valid only under these climatic conditions and in this specific watershed. Yet the study was based on a design of experiments defined on the field expertise and finally led to a full factorial design. This implied a high computational cost due to the large number of simulations (more than 50 000), which does not allow extending systematically it on other climatic conditions and other watersheds.

The present study aims at making possible to represent new watersheds and to determine their sensitivity to input parameters in other climatic and agronomic conditions by reducing the computational cost of the modeling toolkit. A metamodel (or surrogate model) developed on local conditions would allow to perform GSA with low cost yet ensuring it is based on physics. It would help users understanding the most important processes in the VFS they want to design and would increase the operational scope of the modeling toolkit. We performed a much smaller sampling (between 100 and 1000 combinations) of input parameters using the Latin Hypercube Sampling method optimized with a maximin criteria. The metamodel will then be based on a Kriging approach (Rasmussen and Williams, 2006). We will particularly focus on two methodological questions: (i) Is Kriging able to deal with qualitative variables (such as type of soil or type of climate)? (ii) How to integrate the dynamical aspect of state variables in the surrogate construction? The Kriging performance will be analyzed by comparison with other metamodeling methods (Random forest, GAM, linear models; Hastie et al, 2009) and validated with the physically-based simulations conducted on a full factorial test design. Once the metamodel validated, we will assess how the metamodel is a relevant tool to compute global sensitivity analysis in new conditions. The use of the metamodel allows running a large sample of simulations at low cost, such as on Figure 2, which represents the average efficiency and uncertainty of large sample simulations (10 000) by type of soil (from clay-loam, poorly permeable, to sandy soil, very permeable). The metamodel shows a very large variability of efficiency per type of soil but also a large uncertainty of the results. We still need to continue testing and compare methods on prediction uncertainty, sensitivity of prediction quality to the sampling size, and its capability to deal with qualitative variables. Finally, the metamodel will be encapsulated in a user-friendly web interface and will allow comparison of field scenarios, and validation/improvement of actual existing VFS placements and sizing.



Asmussen, L.E., White, A.W., Hauser, E.W., Sheridan, J.M., 1977. Reduction of 2,4-D Load in Surface Runoff Down a Grassed Waterway. *Journal of Environment Quality* 6, 159.

Carluer, N., Noll, D., Bernard, K., Fontaine, A., Lauvernet, C., 2014. Dimensionner les zones tampons enherbées et boisées pour réduire le transfert hydrique des produits phytosanitaires. *Techniques, Sciences et Méthodes* 12, 101-120.

Dosskey, M.G., Helmers, M.J., Eisenhauer, D.E., 2011. A design aid for sizing filter strips using buffer area ratio. *Journal of Soil and Water Conservation* 66, 29–39.

Hastie, Tibshirani and Friedman, 2009. *The Elements of Statistical Learning* (2nd ed.). Springer-Verlag.

Lauvernet C., Muñoz-Carpena R., Carluer N. (2015) Metamodeling as a tool to size vegetative filter strips for surface runoff pollution control in European watersheds. *Geoph.Res. Abstr.* Vol. 17, EGU2015-14672-1.

Muñoz-Carpena, R., J.E. Parsons, et J.W. Gilliam, 1999. Modeling hydrology and sediment transport in vegetative filter strips. *Journal of Hydrology* 214:111.

Rasmussen, C.E. and Williams, C.K.I. *Gaussian Processes for Machine Learning*. The MIT Press, 2006. ISBN 0-262-18253-X.

USDA-SCS, 1972. *National Engineering Handbook, Part 630 Hydrology*. Washington, D.C.

The use of computed assisted semen analysis (CASA) as a method for environmental and toxicological risk assessment – The use of different chambers as sensitivity factor

P. Massanyi¹, N. Lukac¹, R. Stawarz², and J. Danko³

¹Slovak University of Agriculture in Nitra, Nitra, Slovak Republic

²Pedagogical University, Krakow, Poland

³University of Veterinary Medicine and Pharmacy, Kosice, Slovak Republic

Abstract

Reproduction is the biological process by which new individual organisms are produced. It is a fundamental feature of all known life; each individual organism exists as the result of reproduction. Spermatozoa production results in the daily formation of many millions of spermatozoa. The purpose of spermatogenesis is to establish and maintain daily output of fully differentiated spermatozoa that in mammals ranges from more than 200 million in man to 2-3 billion in bull.

Environmental pollution is increasing by rapid leaps worldwide due to the development of modern human society. After industrial revolution, a huge amount of toxic chemicals are disposed into the environment from various anthropogenic activities including industry, agriculture, mines, transportation and settlement. Stress to toxic metals is one of the best examples of evolution driven factors derived by anthropogenic activities. A rapid rate of metal pollution can be a strong selection force causing rapid evolutionary changes in organisms manifested as metal tolerance occurring over time scales as centuries and even decades. Conditions also develop so as to promote uneven distribution of essential elements in the animal organism and change their interaction.

The aim of this study is to describe CASA method useful for estimation of changes related to environmental biology and ecology. In relation to conference topics target of this study is related mainly to modelling (the use of different chambers) related to sensitivity. Semen and spermatozoa analysis evaluates certain characteristics of semen and spermatozoa contained therein. It is done to help evaluate male fertility. Depending on the measurement method, just a various characteristics may be evaluated. The most common reasons for semen analysis are related to infertility investigation.

Motility is the basic parameter used for CASA analysis. Usually in each sample the following parameters were evaluated - percentage of motile spermatozoa (motility $> 5 \mu\text{m/s}$), percentage of progressive motile spermatozoa (motility $> 20 \mu\text{m/s}$), DCL (distance curved line; μm), DAP (distance average path, μm), DSL (distance straight line, μm), VCL (velocity curved line, $\mu\text{m/s}$), VAP (velocity average path, $\mu\text{m/s}$), VSL (velocity straight line, $\mu\text{m/s}$), ALH (amplitude of lateral head displacement, μm) and BCF (beat cross frequency, Hz). From these parameters also others are calculated as linearity - VSL:VCL, straightness - VSL:VAP and wobble - VAP:VCL.

In this study rabbit spermatozoa motility parameters, measured using different evaluation chambers, were compared. The measurement was done using CASA (Computer Assisted Semen Analysis) system; each sample was placed into four different chambers - microscopic slide, Zander Spermometer, Standard Count Analysis Chamber Leica 20 micron and Makler Counting Chamber. CASA showed

that all measured parameters varied depending on chamber used as follows: an average spermatozoa concentration was $1.02 - 1.17 \times 10^6/\text{ml}$, the percentage of motile spermatozoa was in range 59.85 - 77.78% and spermatozoa with progressive motility was ranged from 46.14 to 68.57%. Of other parameters, DAP was $19.23 - 24.44 \mu\text{m}$, DCL $37.43 - 47.20 \mu\text{m}$, DSL $14.27 - 18.92 \mu\text{m}$, VAP $45.26 - 57.31 \mu\text{m/s}$, VCL $87.45 - 110.37 \mu\text{m/s}$, VSL was $33.77 - 44.31 \mu\text{m/s}$, straightness 0.71 - 0.76, linearity 0.36 - 0.40, wobble 0.50 - 0.52, ALH $4.18 - 4.60 \mu\text{m}$ and BCF 23.58 - 28.16. Statistical analysis detected significant differences in almost all studied parameters in regards to evaluation

chamber used. Particularly, highest values for concentration, percentage of motile and progressive motile spermatozoa were detected when microscopic slide with coverslip was used as a chamber. In parameters of the distance, velocity, linearity, straightness and BCF the highest values were obtained using Zander Spermometer, whilst the amplitude of lateral head displacement was the highest in the Makler chamber. These results clearly suggest that the type of evaluation chamber may influence a reliability of measurement of spermatozoa parameters. Our previous in vitro experiments detected various significant dose and time dependent decrease of percentage of motile spermatozoa. Similar tendencies were observed for progressive spermatozoa motility. Detail motility analysis (curvilinear path, average path, straight motility path) show significant differences mainly after various time periods of culture. For further analysis we suggest to calculate concentration dependent decrease of spermatozoa progressive motility up to 50% of control (CDSM50) which should be calculated from at least five replicates using standard statistical tests and compared to progressive spermatozoa motility in control in various time periods.

Financial support: VEGA 1/0857/14, 1/0760/15, APVV-0304-12, KEGA 006/SPU-4/2015

The role of Rosenblatt transformation in global sensitivity analysis of models with dependent inputs

Thierry Mara^{*1}

¹PIMENT, Université de La Réunion, 15 Avenue René Cassin, BP 7151, 97715 Moufia, La Réunion, France

Abstract

Reliable methods exist to perform global sensitivity analysis (GSA) of computer models with independent input factors. In that case, there are different reliable importance measures proposed in the literature to perform this task (e.g. [1–4]). Performing GSA of models with dependent input factors is more challenging. Several has been recently introduced to perform such an analysis ([5–7]). In the cited papers, the authors introduced defined variance-based sensitivity indices for models with dependent inputs and also proposed different methods to assess them.

Following the original idea of [8], the new variance-based sensitivity indices allow for distinguishing the contributions of an input factor to the model response variance that account for its dependence with the other inputs and that do not account for its joint dependence contribution. Such a kind of distinction also stands if one consider another type of importance measure. For instance, in [9], the authors defined these kind of indices for the moment-independent measure of Borgonovo [3]. The aim of my presentation is to highlight the role of Rosenblatt transformation (RT, [10]) for analysing computer models with dependent input factors.

Let $y = f(\mathbf{x})$ be the model response of a computer model function of n random input factors $\mathbf{x} = (x_1, \dots, x_n) \sim p(\mathbf{x})$. RT transforms $\mathbf{x}$ into a random vector $\mathbf{u}$ uniformly distributed over the unit hypercube $\mathbb{K}_n$. RT can be written as follows,

$$\begin{cases} u_1 = F_1(x_1) \\ u_2 = F_{2|1}(x_2|x_1) \\ \vdots \\ u_n = F_{n|\sim n}(x_n|\mathbf{x}_{\sim n}) \end{cases} \quad (1)$$

where F_1 , $F_{i_k|i_1\dots i_s}$ are respectively the cumulative distribution function of $x_1 \sim p_1(x_1)$ (i.e. $p_1(x_1) = dF_1/dx_1$) and the conditional cumulative distribution function of $x_2 \sim p_{2|1}(x_2|x_1)$. The vector $\mathbf{x}_{\sim n}$ stands for $\mathbf{x}/x_n$.

Once the input vector $\mathbf{u}$ obtained, it is straightforward to compute any desired sensitivity index related to the u -variables since they are independent. The interpretation of the sensitivity of $\mathbf{u}$ as those of $\mathbf{x}$ where given first in [6] (actually, during the VI-th SAMO Conference 2010 in Milan). The sensitivity indices of u_1 are those of x_1 , the sensitivity indices of u_2 are those of x_2 without its mutual influence due to its dependence with x_1 , etc. Finally, we note that the influence of u_n onto the model response is simply the one of x_n without accounting for its dependence with the other variables.

^{*}Corresponding author: mara@univ-reunion.fr

To assess the independent influence of all the inputs onto the model response, one must notice that the Rosenblatt transformation is not unique. Indeed, one can start with any input variable x_i in (1) and also ends the transformation with any x -variable. The main drawback of RT is that the conditional densities must be known in advance. But, the advantage is that one can perform GSA with any importance measure and also any method proposed in the literature. Thus, one can adapt the Morris method to the screening of computer models with dependent inputs as proposed proposed by some authors (paper under evaluation).

References

- [1] I. M. Sobol'. Global sensitivity indices for nonlinear mathematical models and their monte carlo estimates. *Mathematics and Computers in Simulation*, 55:271–280, 2001.
- [2] T. Homma and A. Saltelli. Importance measures in global sensitivity analysis of nonlinear models. *Reliability Engineering and System Safety*, 52:1–17, 1996.
- [3] E. Borgonovo. Measuring uncertainty importance: Investigation and comparison of alternative approaches. *Risk Analysis*, 26(5):1349–1361, 2006.
- [4] I. M. Sobol' and S. Kucherenko. A new derivative based importance criterion for groups of variables and its link with the global sensitivity indices. *Computer Physics Communications*, 181:1212–1217, 2010.
- [5] S. Kucherenko, S. Tarantola, and P. Annoni. Estimation of global sensitivity indices for models with dependent variables. *Computer Physics Communications*, 183:937–946, 2012.
- [6] T. A. Mara and S. Tarantola. Variance-based sensitivity indices for models with dependent inputs. *Reliability Engineering and System Safety*, 107:115–121, 2012.
- [7] T. A. Mara, S. Tarantola, and P. Annoni. Non-parametric methods for global sensitivity analysis of model output with dependent inputs. *Environmental Modelling and Software*, 72:173–183, 2015.
- [8] C. Xu and G. Z. Gertner. Uncertainty and sensitivity analysis for models with correlated parameters. *Reliability Engineering and System Safety*, 93:1563–1573, 2008b.
- [9] C. Zhou, Z. Lu, L. Zhang, and J. Hu. Moment independent sensitivity analysis with correlations. *Applied Mathematical Modelling*, 38:4885–4896, 2014.
- [10] M. Rosenblatt. Remarks on the multivariate transformation. *Annals of Mathematics and Statistics*, 43:470–472, 1952.

PC Expansion for Global Sensitivity Analysis of non-smooth functionals of uncertain Stochastic Differential Equation solutions

M. Navarro¹, O.P. Le Maître^{1,2}, O.M. Knio^{1,3}

¹CEMSE Division, King Abdullah University of Science and Technology, Thuwal, KSA.

²LIMSI-CNRS, UPR-3251, Orsay, France.

³ Department of Mechanical Engineering and Materials Science, Duke University, Durham, NC, USA.

Stochastic differential equations (SDEs) play an important role in modeling problems in many different fields. These models are particularly challenging since, in addition to their inherent random dynamics, they are often uncertain due to incomplete knowledge of the parameters and input data, etc. As a result, the model outputs are also uncertain with a variability depending on both the inherent noise and the model uncertainties. The aim of the present work is then to develop efficient techniques to carry out a global sensitivity analysis (GSA) in uncertain SDE driven by Wiener noise. The objective is to quantify the respective contributions, to the total variance of the SDE solution, of the Wiener noise and other sources of parametric uncertainty.

The simplest approach to perform such decomposition of the variance is through a Monte Carlo sampling or any of its variants [2]. Although the implementation of these methods is straightforward, neither MC nor any of its improved versions exploit the potential smoothness with respect to uncertain parameter of the model output, to accelerate the convergence of the sensitivity indices. We propose to use functional approximations (Polynomial Chaos) to account for the dependences on the uncertain model parameters of the random trajectories, as introduced in [1]. Under the assumption that the driving Wiener noise and the uncertain parameters are independent random quantities, the PC expansion can be exploited to perform an orthogonal decomposition of the variance, separating contributions from the uncertain parameters, the Wiener noise, and a coupled contribution. The approach in [1] relied on Galerkin methods to compute the stochastic PC modes of the solution and on the Sobol-Hoeffding decomposition to define the sensitivity indices [5], although others methods such as Fourier amplitude sensitivity test (FAST) can be applied [3]. In the present work, we propose an extension to non-intrusive or sampling methods for the determination of the PC approximation, along with techniques to perform a GSA of any functional of the SDE solution. In more details, we rely here on a non-intrusive pseudo-spectral projection (PSP) method, over a sparse-grid of parameter points, for the PC modes computation. The GSA of QoIs derived from SDE solution is illustrated in Figure 1. Two different QoIs are considered: the case of path integral which inherit the smoothness of the original SDE solution (see *Left*) and the case of exit time which exhibit non smooth dependences with the uncertain parameters (see *Right*).

For the first case, the non-intrusive projection of the QoI can be directly performed, and the sensitivity indices can be computed from the random PC modes. Numerical experiments show that the standard error (SE) of the sensitivity indices for the direct approach can be much less than for the SE of the classical MC estimation [4] of the sensitivity indices, when the same number of Wiener noise samples are used. This improvement in the SE comes at the cost of having to solve as many SDEs as sparse grid points in the parameter space; so the reduction of the variance in the sensitivity indices estimators may not necessarily translates into an improvement of the overall computational efficiency, in particular in the case of complex parametric dependencies demanding large sparse grids.

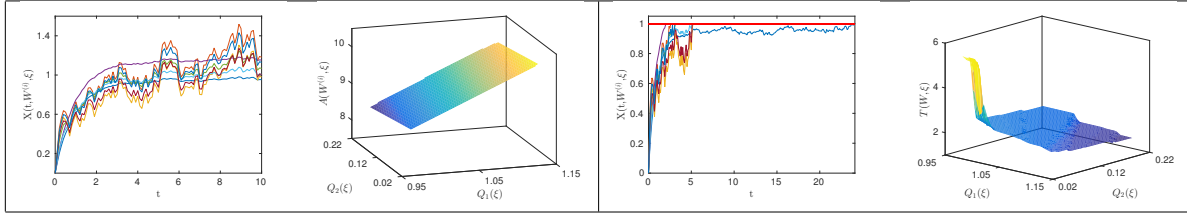


Figure 1: Ornstein-Uhlenbeck (OU) process with uncertain parameters: $dX(W, \xi) = (Q_1(\xi) - X(W, \xi))dt + (\nu X(W, \xi) + 1)Q_2(\xi)dW$, $X(t = 0) = 0$, with parameters $Q_1 \sim \mathcal{U}([0.95, 1.15])$, $Q_2 \sim \mathcal{U}([0.02, 0.22])$ and $\nu = 0.2$. *Left*: Integral over time of X (smooth QoI). *Right*: Exit time of X at level 1 (non-smooth QoI).

For the second case, the QoI (exit time) does not inherit the smoothness of the SDE solution with respect to the uncertain parameters, and its direct non-intrusive projection has a very slow convergence. Therefore, we introduce the idea of *indirect* non-intrusive projection which consists in, first, constructing the approximation of the SDE solution and, second, sampling this approximation to generate conditional samples of the QoI. This approach yields a lower standard error than the classical MC estimator when the same number of noise samples is used, as can be seen in Figure 2. However, the cost of the improvement might again be significant for some problems. Our numerical experiments show that for problems with a low noise effects and low number of sparse grid points, the indirect projection is more efficient than direct MC, providing lower SE for the same cost. In addition, for the indirect projection gives access to a richer information than direct MC sampling (for instance all sensitivity indices can be computed without requiring more noise samples).

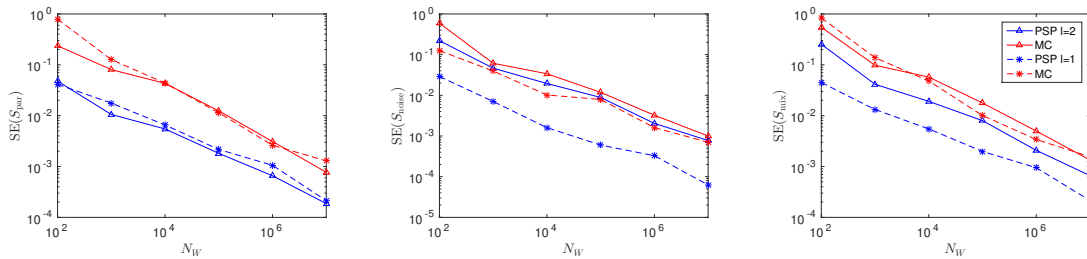


Figure 2: Standard errors in the sensitivity indices of the exit time, as a function of the number of samples N_W , for the indirect PSP and standard MC approaches. The solid lines correspond to the multiplicative noise OU process with large noise, while the dashed lines correspond to the low additive noise case.

References

- [1] O. P. Le Maître and O. M. Knio. Pc analysis of stochastic differential equations driven by wiener noise. *Reliability Engineering & System Safety*, 135:107–124, 2015.
- [2] J. S. Liu. *Monte Carlo strategies in scientific computing*. Springer Series in Statistics, 2004.
- [3] Clémentine Prieur and Stefano Tarantola. *Variance-Based Sensitivity Analysis: Theory and Estimation Algorithms*, pages 1–23. Springer International Publishing, 2016.
- [4] A. Saltelli. Making best use of model evaluations to compute sensitivity indices. *Computer Physics Communications*, 145(2):280–297, 2002.
- [5] I. M. Sobol’. Sensitivity estimates for nonlinear mathematical models. *Math. Modeling & Comput. Exp.*, 1:407–414, 1993.

Statistical emulation as a tool for analysing complex multiscale stochastic biological model outputs

O. Oyebamiji¹ and D.J. Wilkinson¹

¹School of Mathematics & Statistics, Newcastle University, NE1 7RU

Abstract

The performance of credible simulations in open engineered biological frameworks is an important step for practical application of scientific knowledge to solve real-world problems and enhance our ability to make novel discoveries. Therefore, maximising our potential to explore the range of solutions at frontier level could reduce the potential risk of failure on a large scale. One primary application of this type of knowledge is in the management of wastewater treatment systems. Efficient optimisation of wastewater treatment plant focuses on aggregate outcomes of individual particle-level processes. One of the crucial aspects of engineering biology approach in wastewater treatment study is to run a high complex simulation of biological particles. This type of model can scale from one level to another and can also be used to study how to manage real systems effectively with minimal physical experimentation.

To identify crucial features and model water treatment plants on a large scale, there is a need to understand the interactions of microbes at fine resolution using models that could provide the best available representation of micro scale responses. The challenge then becomes how we can transfer this small-scale information to the macroscale process in a computationally efficient and sufficiently accurate way. It has been established that the macro scale characteristics of wastewater treatment plants are the consequences of microscale features of a vast number of individual particles that produce the community of such bacterial (Ofiteru et al. 2014).

Nevertheless, simulation of open biological systems is challenging because they often involve a large number of bacteria that ranges from order 10^{12} to 10^{18} individual particles and are physically complex. The models are computationally expensive and due to computing constraints, limited sets of scenarios are often possible. A simplified approach to this problem is to use a statistical approximation of the simulation ensembles derived from the complex models which will help in reducing the computational burden. Our aim is to build a cheaper surrogate of the Individual- based (IB) model simulation of biological particle. The main issue we address is to highlight the strategy for emulating high-level summaries from the IB model simulation data.

Our approach is to condense the massive, long time series outputs of particles of various species by spatially aggregating to produce the most relevant outputs in the form of floc and biofilms aggregates. The data compression has the benefit of suppressing or reducing some of the nonlinear response features, simplifying the construction of the emulator. Some of the most interesting properties at the mesoscale level like the size, shape, and structure of biofilms and flocs are characterised, see Figure 1. For instance, we characterize the floc size using an equivalent diameter. This strategy enables us to treat the flocs as a ball of a sphere and or fractal depending on the shape, and we approximate the diameter of a sphere that circumscribes its boundary or outline.

We use Gaussian process emulation in the form of kriging metamodels where output data can be decomposed into a mixture of deterministic (non-random trend) and a residual random variation. In particular, we develop dynamic emulators for the multi- outputs simulation data using a multivariate kriging. The kriging model is formulated appropriately to filter the noise derived from replicate simulations. Due to the nature of output data from the simulation model, we use a dynamic emulation technique. Dynamic emulation models the evolution or trajectory of random variables over some time-steps (Conti et al. 2009). Finally, we perform the sensitivity analysis of the kriging model by calculating the total effects of each explanatory variable which helps to identify the relative importance of variables in the model.

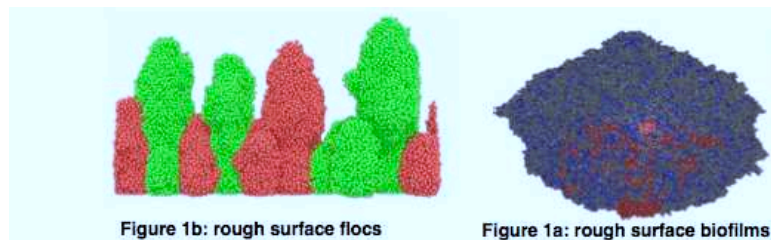


Figure (1) is the simulation data showing the transformation of microscale particles to biofilms and floc aggregates at the mesoscale for a particular time.

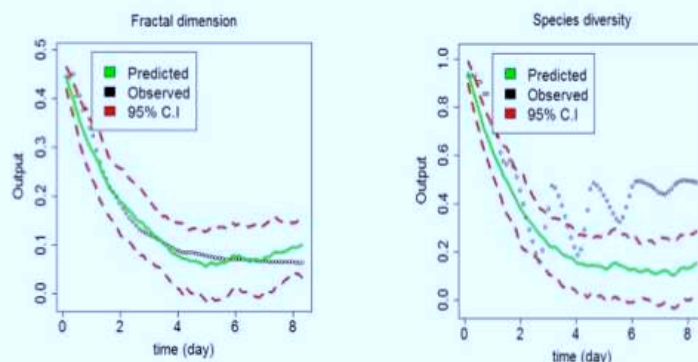


Figure 2: Comparison of the emulator performance with simulation data for two characterized outputs from IB model of floc simulation (black) and their emulator predictions (green) with 95% C.I (red).

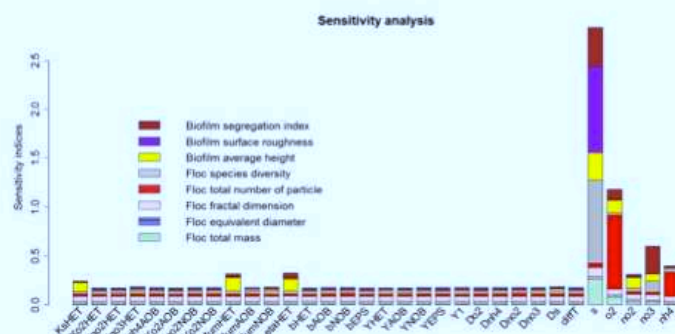


Figure 3: Barplots showing the kriging based sensitivity indices for the eight characterized outputs from IB model.

In Figure (2), the emulator for the fractal dimension predicts the temporal behaviour relatively well, almost all the points lie within the 95% C.I. The predicted bands remain very small. The species diversity emulator produces similar pattern to the simulation data although after day "3" the emulator deviates from the usual trend but not significantly. We note that the shape, size and structure of biofilms and floc are essential operation parameters in the management of wastewater.

The sensitivity analysis in Figure (3) indicates that nutrient boundary conditions are the most critical parameters for predictions of most of the outputs. These parameters regulate the distribution and transports of nutrients across the computational domain thus determine the particle growth and division.

References:

- Conti, S., Gosling, J. P., Oakley, J. E., & O'hagan, A. (2009), Gaussian process emulation of dynamic computer codes. *Biometrika*, asp028.
- Ofiteru, I. D., Bellucci, M., Picioreanu, C., Lavric, V., & Curtis, T. P. (2014), Multi-scale modelling of bioreactor-separator system for wastewater treatment with two-dimensional activated sludge floc dynamics. *Water Research*, 50, 382-395.

Identification of influential parameters in building energy simulation and life cycle assessment

M. L. Pannier¹, P. Schalbart¹, and B. Peuportier¹

¹MINES ParisTech PSL - CES (Center for Energy Efficiency of Systems), 60 Bd Saint Michel 75272 Paris Cedex 06 - France

Abstract

Introduction

Worldwide, the building sector is responsible for high environmental impacts which can be reduced by applying an eco-design approach in new construction and renovation. This requires dedicated tools for the dynamic building energy simulation (DBES) and the building life cycle assessment (LCA). However, the input variables or parameters of these models are uncertain and induce variability. Therefore, the reliability of the output has to be investigated, in order to make robust decisions in an uncertain context. A local sensitivity analysis method, Morris screening method and a global sensitivity analysis method based on the calculation of Sobol' indices were applied and their results compared. Using *Pléiades+COMFIE* and *novaEQUER*, respectively for the DBES and LCA, the most influential parameters were identified, and the tools' robustness was tested.

Model description

In *Pléiades+COMFIE*, the building is divided into thermal zones with homogeneous temperature (Peuportier and Blanc-Sommereux, 1990). The buildings elements are meshed and energy equations are solved to obtain the temperature, heating and cooling loads. Twelve environmental indicators are calculated using *novaEQUER* (Polster, 1995; Popovici, 2005). Energy load and geometry data are transferred from *Pléiades+COMFIE*. For the use phase, which is predominant due to the long building lifetime, information about water consumption, occupants transportation, and waste production complement the energy loads. The LCA database ecoinvent provides inventories and impact assessment indicators related to these processes, and the fabrication and end of life building materials.

The model has a large number of equations that constitute a time-dependent non-linear system. Despite its physical nature, it is not possible to easily relate the outputs to the inputs. The computing time is decreased by reducing the order of the model. It ranges from a few seconds for a small house with few thermal zones, to a few minutes for an entire and complex district.

Application of sensitivity analysis methods

Pléiades+COMFIE was validated by software inter-comparisons (e.g. IEA's BESTEST) and against experimental data (Munaretto, 2014; Recht et al., 2014).

The robustness of building LCA results is also important because such tools aim at guiding the decision-making process towards more sustainable built environment. Before investigating the potential change in the ranking of several design alternatives, the influence of the uncertain factors was studied by applying and comparing three global sensitivity analysis (GSA) methods having different computation time versus precision compromise (Pannier et al., 2016). 22 uncertain factors from both energy and environmental models were investigated with an adapted local sensitivity analysis (ALSA), a Morris screening and a GSA. In this case study (single passive family house), the computation time varied between 4 min and 18 h depending on the method.

For the ALSA the difference between the model outputs at each boundary of the variation range was calculated. The indices were divided by the sum of all indices to get the factor's influence, instead of the sensitivity of the model to this factor as usually calculated with LSA. For the Morris screening, 6 levels and 50 OAT repetitions were chosen and the Euclidian distance to the Morris graph origin d_j^* was calculated as in (1) to rank the uncertain parameters according to both the linear effects and the non-linear or interactions effects:

$$d_j^* = \sqrt{\mu_j^{*2} + \sigma_j^2} \quad (1)$$

with μ_j^* the mean of the absolute value of the elementary effect for the j -th factor and σ_j the standard deviation of the elementary effect. Lastly, truncated normal distributions and 1000 samplings were chosen for almost all factors for the GSA. The ranking is based on Sobol' indices.

All three methods identify the same uncertain factors to be the most influential. In all cases, the type of the electricity production mix, the building lifetime, and some factors influencing the energy performance, are the drivers for almost all environmental indicators. However, the relative influence of the factors for all methods is different and the factor ranking change. This is due to methodological differences. The ALSA does not catch the interactions between parameters or the non-linearity. In Morris screening, using a regular grid is equivalent to making the assumption of a uniform distribution for the factors. However, truncated normal distributions were chosen in GSA. So the range boundaries are more explored than with the GSA. Lastly, ALSA and Morris screening evaluate the effect of the factors on the outputs whereas the GSA calculates the variances.

The choice of the most adapted method depends of the scope of the study and the knowledge of the studied building. If the building is well known and small uncertainty ranges are defined, an ALSA can be sufficient. Furthermore, if the factors' sensitivity must be precisely known, a Morris screening can be performed before a GSA in order to reduce the computation time.

Conclusion

The sensitivity analysis methods presented in this paper were applied in a DBES and a building LCA tool in order to identify the most influential factors of the models in the frame of model validation and robustness analysis. It can also be extended to optimisation (Recht et al., 2016) or model calibration as in Robillart (2015) where a Morris screening was used before an approximate Bayesian computation.

Acknowledgements: These projects were performed in the frame of the research project ANR FIABILITE and the research Chair ParisTech "Eco-design of buildings and infrastructure" supported by VINCI.

References

- Munaretto, F. (2014). Etude de l'influence de l'inertie thermique sur les performances énergétiques des bâtiments, PhD Thesis, Ecole Nationale Supérieure des Mines de Paris, 322 p.
- Pannier, M.-L., Schallbart, P., and Peuportier, B. (2016). Identification de paramètres incertains influents en analyse de cycle de vie des bâtiments. In IBPSA France (Champs-sur-Marne), 8 p.
- Peuportier, B., and Blanc-Sommereux, I. (1990). Simulation tool with its expert interface for the thermal design of multizone buildings. *Int. J. Sol. Energy*, 8, 109-120.
- Polster, B. (1995). Contribution à l'étude de l'impact environnemental des bâtiments par analyse du cycle de vie, PhD Thesis, École nationale supérieure des mines de Paris.
- Popovici, E. (2005). Contribution to the life cycle assessment of settlements, PhD Thesis, 244 p.
- Recht, T., Munaretto, F., Schallbart, P., and Peuportier, B. (2014). Analyse de la fiabilité de COMFIE par comparaison à des mesures. Application à un bâtiment passif. In IBPSA France (Arras), 8 p.
- Recht, T., Schallbart, P., and Peuportier, B. (2016). Ecodesign of a "plus-energy" house using stochastic occupancy model, life-cycle assessment and multi-objective optimisation. In *Building Simulation & Optimisation*, (Great North Museum, Newcastle: To be published), 8 p.
- Robillart, M. (2015). Etude de stratégies de gestion en temps réel pour des bâtiments énergétiquement performants, PhD Thesis, École nationale supérieure des Mines de Paris, 258 p.

Sensitivity analysis as essential tool to gain insight into potential hydrological change due to coal development in Australia

Luk Peeters¹

¹CSIRO Land and Water, Australia

Abstract

The development of coal resources through mining or through extracting coal bed methane will potentially affect water resources. Coal bed methane extraction requires the depressurisation of coal seams at depth, the effects of which can propagate through to shallower aquifers and change groundwater levels or surface water groundwater interaction fluxes. Coal mining also has a direct impact on groundwater as coal seams are dewatered for mining. In addition to that, open cut mines intercept rainfall across the footprint of their workings, reducing the runoff to streams.

The Bioregional Assessment Programme is a four year research project to evaluate the direct, indirect and cumulative impacts of coal development across parts of eastern Australia. The Programme identified six bioregions, subdivided into 13 subregions, in which there is the potential for coal resource development. One of the goals of the research is to provide a probabilistic estimate of the change caused by the most likely resource development pathway on water dependent assets in each subregion. Each asset is linked to a set of hydrological response variables, summaries of the hydrology relevant to an asset. Examples of those are the maximum change in groundwater level at an irrigation bore or the change in the number of low flow events in a stream.

The goal of estimating the change in hydrology probabilistically is to fully capture the model predictive uncertainty so as to inform a risk based management of the water resources. This starts with developing a chain of models that is able to numerically simulate the difference in each hydrological response variable between a baseline future and a future with the most likely coal resource developments included. The posterior predictive probability distribution for each hydrological response variable at each location is obtained by integrating the model chain into an approximate Bayesian computation Markov Chain Monte Carlo framework (ABC-MCMC) to constraint the prior parameter distributions with the available relevant observations. While not formally based on the Bayesian likelihood functions, the ABC is preferred as it allows expert knowledge to be included explicitly in the analysis in the absence of sufficient data to establish proper error models for the available observations. In practice, it relies on experts specifying a set of criteria for which model results are deemed acceptable.

The probabilistic, numerical evaluation of predictive uncertainty can however only capture a part of the uncertainty as each numerical model has a set of inbuilt assumptions and model choices that are not straight forward to include in a probabilistic uncertainty analysis. The assessment is therefore complemented by a structured discussion and justification of the main assumptions and model choices in terms of the limiting factors necessitating the assumption (data, resources, technical) and the perceived effect on the predictions.

For the Markov Chain Monte Carlo process, the entire model chain needs to be evaluated 100s if not

1000s of times, which represents a vast computational burden. In addition to that, creating a robust computational framework to integrate a variety of models and run them in sequence is operationally very challenging. For these reasons, in the Markov Chain Monte Carlo process, the original model is replaced with a Gaussian Process emulator. The design of experiment for the training of the emulator is based on a dense Latin hypercube sampling of parameter space. Such emulator can be created for each hydrological response variable at each location very quickly. This allows to tailor posterior parameter distributions to individual predictions with the ABC MCMC.

The parameterisation of the chain of models invariably leads to parameters that have little or no effect on a particular prediction. For this reason, a sensitivity analysis of the design of experiment results is a routine part of the modelling protocol. The main goal is to have a structured procedure in place to guide the prioritisation of factors for inclusion in the Gaussian Process emulators. In addition to that, the sensitivity analysis enables us to focus attention in defining and eliciting prior distributions for parameters. The sensitivity analysis uses the density based sensitivity metrics introduced in Plischke et al (2013). These metrics are augmented with scatter and frequency plots of parameter values versus prediction for selected predictions. These plots both serve as a reality check and to communicate the procedure.

In one of the regions the sensitivity analysis was able to show that both faulting and surface water groundwater interaction, both processes a priori considered to be highly influential on the predictions, were less important than having information on the hydraulic properties of the stressed aquifer or the way the hydrological characteristics of the landscape were captured in the numerical models.

In another region the sensitivity analysis highlighted that the change in groundwater level is mostly affected by the vertical hydraulic conductivity, while the available groundwater level observations to constrain the model were only sensitive to changes in recharge and river bed conductance.

Routinely applying a robust, global sensitivity analysis to the design of experiments proved to provide invaluable insights in the workings of the numerical models and the underlying physical systems. The added value of this understanding is that it, in addition to probabilistically estimating the hydrological change for a specific coal resource development pathway, provides clear guidance for future model development, data collection and monitoring.

References

Plischke E, Borgonovo E, and Smith CL (2013) Global sensitivity measures from given data, *European Journal of Operational Research* 226, 536-550

Global sensitivity analysis with distance correlation and energy statistics

Elmar Plischke¹, Emanuele Borgonovo²

¹ Institut für Endlagerforschung, TU Clausthal

² Department of Decision Sciences, Univ. Bocconi

In scientific modeling, it is often impossible to grasp the response of the output of a numerical model to variations in the model inputs based on sole intuition. For this, sensitivity analysis comes at hand, allowing to analyze the influence of input factors on the model output. Interpreting these sensitivity measures as distances to certainty, one arrives at sensitivity measures based upon tests for stochastic independence. Recently, a number of omnibus tests for this purpose have been suggested in the statistics and machine learning literature, based on distance covariance, energy statistics or the Hilbert-Schmidt independence criterion.

We consider the appropriateness of these measures for sensitivity analysis purposes. In particular, we study the energy statistics. Its one-dimensional analogon is known as Gini mean distance. We embed it into a sensitivity framework recently established by the authors and derive simple estimators so that a cheap method for extracting moment-independent sensitivity information from given data is obtained. Links to reliability theory and to variance-based first order effects can also be established, allowing an interpretation of Gini sensitivity as a mean output quantile sensitivity.

General Framework for Sensitivity Measures

Most of the global sensitivity measures available can be embedded in the following framework which consists of a suitable distance or divergence operator $\zeta(\cdot, \cdot)$ which measure distances or divergences between point estimates, cumulative distribution functions, probabilistic density functions or characteristic functions of Y and of Y conditional to an input parameter $X_i = x_i$ (or a group of input parameters) being set to a specific value. Averaging over all possible input values yields the sensitivity measure

$$s(X_i) = \mathbb{E}[\zeta(Y, Y|X_i)].$$

Moreover, the sensitivity framework can be embedded into a value-of-information context for proper scoring functions, giving rise to . For this, suppose that a decision maker has to select between a set of A alternative strategies, leading to the expected value of perfect information measure

$$\varepsilon_i = \mathbb{E} \left[\max_{a=1, \dots, A} \mathbb{E}[Y^a | X_i] \right] - \max_{a=1, \dots, A} \mathbb{E}[Y^a]$$

Replacing the selection of one of the alternatives with a scoring function we obtain the value-of-information sensitivity

$$\varepsilon_i^S = \mathbb{E} \left[\max_{a=1, \dots, A} \mathbb{E}[S(Y, a) | X_i] \right] - \max_{a=1, \dots, A} \mathbb{E}[S(Y, a)]$$

Proper scoring rules are convex functions which attain their maximum at the value to be estimated (optimal reporting). Hence in this context, the decision is to use the best estimate. For example, consider the Brier score $S(y, a) = -(y - a)^2$ then the value-of-sensitivity measure is the variance of the conditional expectation, the unnormalized version of the variance-based first-order effect. Scoring functions can be used to report point estimates, but may also report distributional forecasts, be it in the form of a density or a distribution. The continuous ranked probability score (CRPS),

$$S(F, a) = - \int (F_Y(y) - \mathbf{1}\{a < y\})^2 dy$$

measures the difference between the predicted and the observed cumulative distribution. CRPS can also be obtained by averaging the so-called quantile/tick/check scores.

Energy statistics, Gini mean difference, Cramér-von Mises distances

The CRPS-induced value-of-information sensitivity indicator is given by

$$\varepsilon_i^{Gini} = \varepsilon_i^{CRPS} = \iint (F_Y(y) - F_{Y|X_i=x}(y))^2 dF_{X_i}(x) dy.$$

Hence, in this case, ζ is a Cramér-von Mises distance. Its multidimensional version called energy statistics is now under heavy investigation. The term $2 \int F(y)(1 - F(y))dy$ is called Gini mean difference and is forming part of the CRPS sensitivity indicator. We therefore call this measure Gini sensitivity. With a scatterplot partition approach, replacing the atomic condition $X_i = x_i$ with an interval condition $X_i \in [F_{X_i}^{-1}(\frac{m-1}{M}), F_{X_i}^{-1}(\frac{m}{M})]$ using a partition into M equipopulated classes one can form the empirical conditional cumulative distributions and use formulas for the energy statistics which are based on Hoeffding U-statistics or a direct estimate. A version of the Gini measure which is invariant with respect to monotonic transformations is obtained with

$$\varepsilon_i^{Gini*} = 6 \iint (F_Y(y) - F_{Y|X_i=x}(y))^2 dF_{X_i}(x) dF_Y(y).$$

The scale factor is chosen in such way that this measure is located in [0,1]. Note that the underlying ζ is now a divergence without symmetry in the arguments. Using a quantile idea we can obtain this as mean of unnormalized first-order effects of the quantile indicator functions,

$$\varepsilon_i^{Gini*} = 6 \int_0^1 \mathbb{V}[\mathbb{E}[\mathbf{1}\{Y \leq F_Y^{-1}(\theta)\} | X_i]] d\theta.$$

Hence, also PCE or other regression methods might be used for obtaining estimates. With respect to reliability theory, this can be considered an average over limit state functions when the limit varies.

We also consider the distance covariance and assess if it is a suitable sensitivity measure.

The use of the Gini sensitivity measure as moment-independent importance measure (MIM) circumvents a few problems with hinders the numerical estimation of other MIMs, e.g. there is neither a pdf to be estimated, nor numerical instable maxima have to be taken into account. Moreover, the square reduces effectively numerical noise in the estimation. We are looking forwards to gain further experience with these measures.

Combining switching factors and filtering operators in GSA to analyze models with climatic inputs

SÉBASTIEN ROUX

INRA, UMR SYSTEM, 34060 Montpellier, France

SAMUEL BUIS

INRA, UMR EMMAH, 84914 Avignon, France

PATRICE LOISEL

INRA, UMR MISTEA, 34060 Montpellier, France

Introduction

This work is devoted to the analysis of models having functional inputs and is motivated by the intensive use of climatic variables in crop models. The main output of these models is the crop yield, which is estimated, among others, from daily-sampled climatic variables (Temperature, Rain, Radiation, Evapotranspiration). We want to test to what extent this fine temporal resolution is mandatory to generate accurate predictions and quantify how much *a priori* simplifications, such as lowering the temporal resolution, would affect the model results. This may lead to a better understanding of model behavior as well as to a simplification of the model and/or of the acquisition of its input variables.

To this aim, we introduce the use of filtering operators into Global Sensitivity Analysis using switching factors [1]. Low pass filters are used to reduce the temporal resolution of climatic variables. GSA is required because we want to explore the impact of this input simplification in a global exploration of model inputs. Switching factors have been proposed [1], [2] in the context of spatially distributed inputs and further analyzed in [3]. They were initially introduced to assess the sensitivity to the presence of stochastic errors in spatial functional inputs. We use them here to test the sensitivity of a model to simplifications of the temporal structure of its climatic inputs.

Methodology

The method is presented for a model f having two independent inputs: one functional (X) and one scalar (p). The scalar output Y is written as $Y = f(X, p)$. Let g be an operator that transforms X into $g(X)$. We introduce a switching factor η and a modified model f_g such as:

$$f_g(\eta, X, p) = \begin{cases} f(X, p) & \text{if } \eta = 0 \\ f(g(X), p) & \text{if } \eta = 1 \end{cases}$$

A map-labeling scheme [4] is used to perform a GSA using n_X samples of X . This leads to a GSA on the independent factors (η, l, p) of model f_g defined as $\tilde{f}_g(\eta, l, p) = f_g(\eta, X_l, p)$, with X_l denoting the sample of X with index l . As we are in a factor fixing context, we focus on Total Sensitivity Indices (TSI) with the perspective of getting small indices for the switching factor η .

Application to a simplified crop model

We tested the method with a simplified crop model that couples a simple water balance with a radiation-driven biomass growth. The crop growth has two limiting factors: High temperature and water stress. The model has the typical input structure of classical crop models (4 climatic variables at a daily time-step). The use of a simplified definition allows for fast computations and qualitative validation of sensitivity results. We applied the proposed method using $n_x = 42$ climatic years and with 3 other factors. TSI of the 4 switching factors, of the climate label and of the 3 others factors were computed using a Sobol algorithm from the R package "Sensitivity". We used two different filtering operators: an inter-annual mean $g_1(X)(t) = \frac{1}{n_x} \sum_{i=1}^{n_x} X_i(t)$ and a local mean operator $g_2^d(X) = X * G_d$, where $*$ denotes the convolution operator and G_d is a mean filter over $(2d + 1)$ days, with d between 1 and 25.

The results presented in Fig. 1 show that the application of the inter-annual filter g_1 does not provide small TSI for both the 4 switching variables ($TSI(\eta_{rain}) > 0.2$). It means that simplification by inter-annual mean is not acceptable. This was however not the case when excluding the rain input from the simplification: In that case, the TSI of the 3 remaining switching factors are simultaneously small. Hence the simplification of these 3 variables by their inter-annual mean seems acceptable.

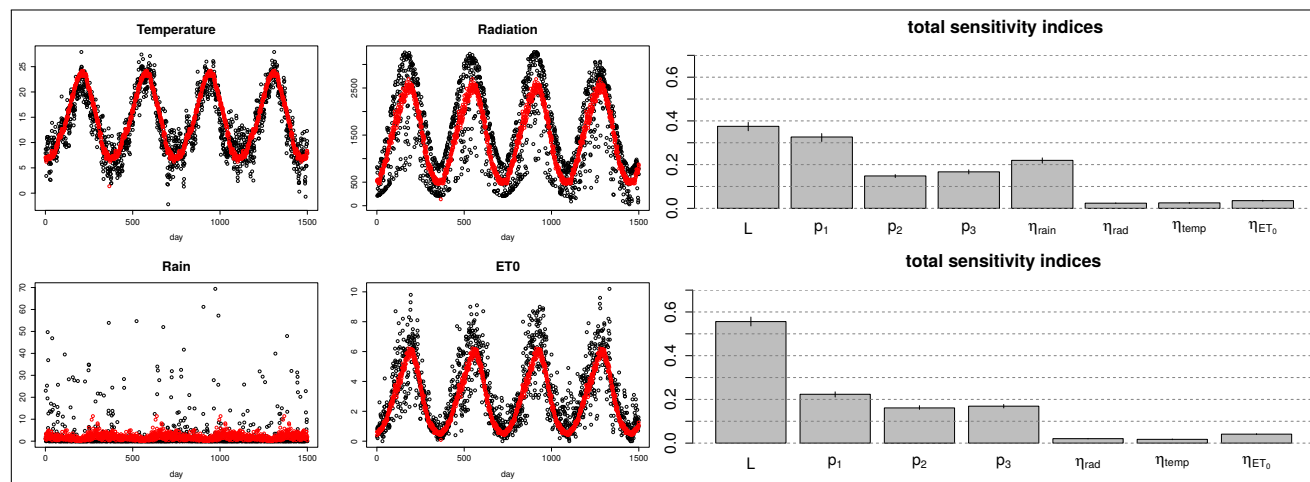


Figure 1: Left: Application of the inter-annual mean operator g_1 (in red) on the climatic variable (year 1..4 in black); Right: Total sensitivity indices when applying the filter to the 4 climatic variables (top) and to 3 of them when excluding Rain (bottom).

Using filter $g_2(d)$ we also showed that rain input can be simplified in terms of temporal resolution even with large d , but only for model settings corresponding to low run-off soil conditions. This property was expected from the model definition.

Conclusion

Combining switching factors with filtering operators to analyze models with climatic inputs seems to be a promising method. This approach can easily be extended to other input transformations in order to test the effect of other *a priori* modifications of the model input structure: An example would be a low pass filtering preserving high value applied to the rain input. One difficulty lies in the definition of such relevant transformation: It requires relevant *a priori* knowledge on the model behavior. A next step of this work is to apply the method to a more complex crop model sharing a comparable structure of model inputs.

References:

- [1] M. Crosetto and S. Tarantola *International Journal of Geographical Information Science*, vol. 15, no. 5, pp. 415–437, 2001.
- [2] S. Tarantola, N. Giglioli, J. Jesinghaus, and A. Saltelli *Stochastic Environmental Research and Risk Assessment*, vol. 16, no. 1, pp. 63–76, 2002.
- [3] B. Iooss and M. Ribatet *Reliability Engineering & System Safety*, vol. 94, no. 7, pp. 1194–1204, 2009.
- [4] L. Lilburne and S. Tarantola *International Journal of Geographical Information Science*, vol. 23, no. 2, pp. 151–168, 2009.

[Sébastien Roux; rouxs@supagro.inra.fr]

RepoSTAR –New Framework for Statistic Runs for Uncertainty and Sensitivity Analysis of a Radioactive Waste Repository Model

T. Reiche¹ and D.-A. Becker¹

¹Gesellschaft fuer Anlagen- und Reaktorsicherheit (GRS) gGmbH, Germany

For long-term assessment of the safety of final repositories for radioactive waste the program package RepoTREND is being developed and applied by GRS [1-2]. The statistical framework of RepoTREND is called RepoSTAR. The development of this framework was motivated by a number of practical problems that occurred in the past when performing probabilistic analyses.

RepoTREND requires a high number of input parameters. A probabilistic investigation should give information about the effects of epistemic or aleatory uncertainties on the safety assessment results. These uncertainties, however, correspond to aspects of the physical problem and not to the technical input requirements of the model. It is possible that one physical uncertainty influences several different program input parameters, maybe even belonging to independent computation modules, and it is also possible that one input parameter is influenced by several principally independent physical uncertainties. In the past, such interconnections often required individual, error-prone programming work that had to be changed for each case study.

To avoid this problem, RepoSTAR, unlike former computational approaches, clearly distinguishes between *probabilistic variables*, which refer to the physical problem, and *input parameters*, which refer to the program. Users are free to define the case study, identify relevant physical uncertainties and parameterise them by defining probabilistic variables, regardless of the model input requirements. Only in the second step, the user has to think about how each input parameter is affected by these variables. RepoSTAR provides a simple formula notation, which allows defining nearly any kind of relationship. By evaluating such dependencies during runtime, but outside the calculation code, this concept offers a convenient practical way to perform probabilistic analyses without any need to modify the code.

A probabilistic analysis with RepoTREND and RepoSTAR proceeds as follows:

Once the problem has been defined, the user has to assign a marginal pdf to each probabilistic variable, according to available expert knowledge. Linear statistical dependencies between the variables can be taken into account by defining a correlation matrix. The user can choose between different sampling schemes (random, LHS, quasi-random, (E)FAST, among others) and random number generators.

In a second step, the links between the probabilistic variables and the program input parameters are defined by the user. These can include mathematical functions as well as logical decisions. The mappings are coded using a simple formula notation, which is fed into the program via the data supply tool XENIA. After defining some general data like the number of runs to be executed, the user starts the probabilistic calculation. For each single run, RepoSTAR evaluates the mapping formulas, replaces the input parameter values accordingly and executes the model run. After completion of each single run, RepoSTAR collects the model output of interest. In principle, any time-dependent model output can be chosen. These values are interpolated to a common time grid and written into a specifically formatted file.

When all runs are finished, the user starts the evaluation tool RepoSUN. This tool is specifically designed to work with the output generated by RepoSTAR. It has an own GUI and provides uncertainty analysis as well as a number of graphical and numerical sensitivity analysis methods (SRRC, PCC, PRCC, Smirnov test, EASI, FAST, EFAST, Sobol, among others).

RepoSUN uses the new SimLab 4 library, which has been developed by JRC and GRS. SimLab 4 contains an interface to the statistical programming language R and can be enabled with low effort to calculate nearly any UA or SA measure. Via SimLab 4, RepoSUN has access to a wide, extendible variety of scripts written in R and is therefore fit for future developments. The RepoSTAR concept can be applied not only in the context of the processes in the final repositories for the radioactive waste but universally, e. g. for any complex model and any computational code.

Acknowledgement: This work was funded by the German Federal Ministry for Economic Affairs and Energy (BMWi) under grant No. 02E770240.

References

- [1] Becker, D.-A.: RepoSTAR - Ein Codepaket zur Steuerung und Auswertung statistischer Rechenläufe mit dem Programmpaket RepoTREND. Gesellschaft für Anlagen - und Reaktor - sicherheit (GRS) gGmbH, GRS-411, BMWi-FKZ 02E10367, Braunschweig, 2016.
- [2] Reiche, T.: RepoTREND - Das Programmpaket zur integrierten Langzeitsicherheitsanalyse von Endlager-systemen. Gesellschaft für Anlagen- und Reaktorsicherheit (GRS) gGmbH, GRS-413, BMWi-FKZ 02E770240, Braunschweig, 2016.

Sensitivity Analysis for Energy Performance Contracting in new buildings

L. Rivalin^{1,2}, P. Stabat¹, D. Marchio¹, M. Caciolo², and F. Hopquin²

¹Mines ParisTech, PSL - Research University, CES - Center for Energy efficient Systems, Paris, France

²ENGIE- Axima - Service Engagement Energétique, Nantes, France

Abstract

Since buildings are responsible for around 40% of total energy, many initiatives have been made to guarantee the savings to the purchaser. Before a building construction, an energy performance contracting (EPC) consists in predicting the energy required to maintain the user's comfort by using thermal dynamic simulation tools. Many building and HVAC system characteristics are uncertain due to lack of knowledge on uncommitted parameters at design stage or implementation defaults at construction stage. To define a performance guarantee, not only a consumption threshold for the performance contracting taking into account uncertainties should be set, it also important to identify the key parameters to pay special attention to during the design phase in order to reduce the risks of non-compliance of the contract. The identification of the most influencing parameters is part of a quality procedure.

This article focuses on several sensitivity analysis methods all useful for an Energy Performance Contracting:

- Quick sensitivity analysis methods to identify the most influencing parameters and to reduce the uncertainty analysis
- Reliability methods to assess part of uncertainty due to the parameters in the failure probability factors for a given consumption threshold
- Global sensitivity methods studying all the input space.

The aim of this paper is to draw recommendations concerning the use of these different methods according to the time budget, the expected accuracy and the type of building parameters to be studied.

The selected methods are assessed on a real case study, a 4000 m² office and stock building of 2 levels located in Nantes (West of France) with 2 air handling units to ensure indoor air quality and 2 reversible Heat Pumps supplying underfloor heating systems, chilled beams and AHUs. The building is modeled under TRNSYS 17 environment. For this building, 49 uncertain parameters of the building and its systems have been identified: AHUs, heat pumps, water networks, building walls, building glazing, infiltration, ground exchanges, set points and occupancy. Three types of probability density are chosen to characterize the parameters, depending on the knowledge we have of the parameters: uniform, beta and truncated normal distributions.

We firstly study two local sensitivity methods: Morris method (screening) and differential sensitivity analysis. Our aim is to reduce the number of input parameters of the probabilistic study by identifying the most influential ones. The ranks obtained by Morris method or Quadratic Combination are very similar, except slight changes between groups, but quadratic combination requires 4 times less calculation than Morris method. The Morris method provides additional information such as the

detection of non-linear effects of the parameters and the interactions between factors. If the physical model contains parameters with very non-smooth effects (for instance, threshold effects correlated with several input parameters), the quadratic combination is not appropriate anymore, so that Morris method is recommended in order to select the most influential input parameters. In the other cases, Quadratic Combination is advised. This first step permits to reduce the number of parameters from 49 to 24.

Then, we study FORM/SORM reliability methods. These methods permit to compute the probability of exceeding a threshold and to assess the contribution of each parameter on the failure probability for a given consumption threshold. These local sensitivity methods are very interesting for EPC. Indeed, the importance factors permit to identify which parameters are critical to maintain a building consumption and thus to adapt the measurement protocol. As the FORM method's cost is low, it can be applied to different consumption thresholds to study the evolution of the share of responsibility of the parameters in the probability of exceeding these thresholds.

Finally, global sensitivity methods are assessed. Since the sampling sensitivity analysis methods cannot be applied directly in our case, an approximation method is used. Indeed, given the computational time of one simulation with our physical model (11'), computing Sobol' indices would require around 24000 simulations to study the selected parameters, that is to say, several weeks of computational time. Thus, as the computational time to apply global sensitivity methods are too high, one solution is to approach the physical model by a much faster model constructed by analysing the effect of the input random variables on the outputs.

The physical model is replaced by a sparse Polynomial Chaos expansion metamodel which runs a building simulation in less than one second. The advantage of the sparse polynomial chaos expansion is that it easily provides Sobol' indices by analysing the coefficients. The first order Sobol' indices identify the most influential model parameters. Sparse chaos polynomial expansion permits to approach a physical model in a very efficient way, with less than 500 simulations, depending on the number of parameters. The approximation of a model by a sparse polynomial chaos works if the model is smooth enough, for example, without threshold effects.

Investigating the Scale Effect of Watershed Delineation on Local Multi-Criteria Method for Land Use Evaluation

Seda Şalap-Ayça, Piotr Jankowski

*San Diego State University – Department of Geography, 5500 Campanile Dr., San Diego, CA, USA,
92182-4493*

Land use evaluation for semi-natural habitats is a land performance process involving careful consideration of several environmental factors to balance the preservation of nature and wildlife while considering the need for agricultural production. Since multiple criteria are involved in the evaluation process, and many areas are subject to high heterogeneity of land characteristics (criteria), local multi-criteria evaluation (MCE) is a suitable approach to solving this decision making problem. As opposed to global MCE, the local approach explicitly accounts for the local variability in evaluation criteria within defined local units of analysis (subregions). Moreover, criteria weights are also assumed non-stationary and are standardized based on the local extreme values within these subregions. In a study of local MCE presented here, watershed delineations are selected for the local unit of analysis, since they reflect the hydrologic principles and are widely applied as a logical unit of land management. The sub-watershed delineations, however, influence how differences in expected decision option outcomes at a large scale (small area) impact the solution of a local model as compared to a global model. Additionally, the quality of the decision support relying on MCE output can be significantly improved by assessing the uncertainty associated with the decision problem. The interaction between the decision model input and output is key information for establishing the level of confidence in the evaluation outcome. Especially, the reliability of outcomes becomes a crucial requirement when the model output is used for decision making affecting environmental sustainability.

An integrated approach to uncertainty and sensitivity analysis (iUSA) can help uncover the sources of uncertainty through the analysis and visualization of input-output relationships captured by sensitivity analysis. This research reports on a study of a local MCE approach followed by uncertainty and

sensitivity analysis of land prioritization model for conservation practices. The relationship between the results of uncertainty-sensitivity analysis and different scales of analysis units is investigated in the local MCE model. The comparison of suitability and uncertainty maps for each scale of land unit is made using the coefficient of variation values per parcel and the relative change in the overall rank. Finally, sensitivity maps are combined into dominance maps by using output variability (first-order) and interaction (total-order) maps. The overview of the implementation is illustrated in Figure 1.

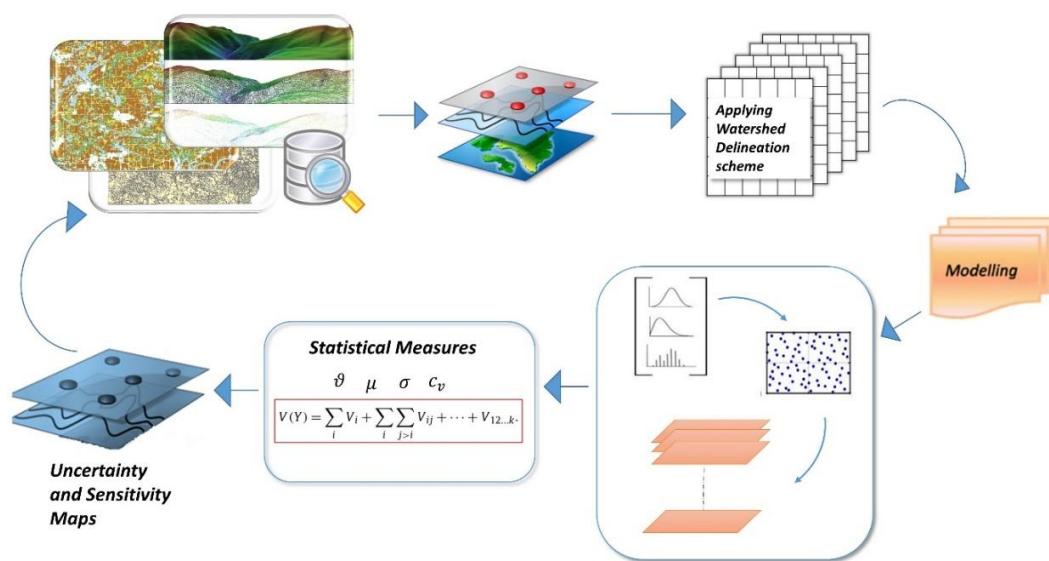


Figure 1. Graphical Abstract of the Framework

The iUSAapproach goes beyond the conventional practice of one-at-a-time (OAT) sensitivity analysis by providing spatial uncertainty and sensitivity maps. Moreover, two different granularity levels of delineations for local MCE used to calculate land parcel suitability show that it is insightful to examine the effect of scale on the stability of MCE model output. A potential practical application of the presentedapproachis the improved analytical support for land suitability evaluation requiring an explicit consideration of multiple decision alternatives.

Keywords: GIS, Local Multi-Criteria Evaluation, Uncertainty Analysis, Global Sensitivity Analysis, Scale Effect

PCE-based Sobol' indices for probability-boxes

R. Schöbi* and B. Sudret*

**ETH Zürich, Institute of Structural Engineering,
Chair of Risk, Safety & Uncertainty Quantification
Stefano-Franscini-Platz 5, CH-8093 Zürich*

April 22, 2016

Abstract

In modern engineering, complex systems and processes are modelled by computational simulation tools, such as finite element models (FEM), to avoid performing expensive physical experiments. Such simulations map a set of input parameters through a computational model to a quantity of interest (QoI). The input parameters are often not perfectly known, because they are estimated based on noisy measurements and expert judgement or because they have an intrinsic variability. This, in turn, introduces uncertainty in the QoI.

The uncertainty in the input parameters are often represented probabilistically by means of joint probability distributions. The distributions are whenever possible inferred from available data. However, a common situation in practice is to have only scarce or incomplete data to characterize a probabilistic random variable. This introduces epistemic uncertainty (lack of data, lack of knowledge) alongside aleatory uncertainty (natural variability) as a source of uncertainty. In such cases, a more general framework, such as probability-boxes (Ferson and Ginzburg, 1996; Oberguggenberger et al., 2009), is required to describe the input uncertainty appropriately. Probability-boxes (p-boxes) describe the cumulative distribution function (CDF) F_X of a random variable X by lower and upper boundary curves, *i.e.* $\underline{F}_X$ and $\overline{F}_X$, respectively. The true but unknown CDF lies between those boundaries, *i.e.* $\underline{F}_X(x) \leq F_X(x) \leq \overline{F}_X(x)$, $\forall x \in X$. A p-box provides a natural framework to distinguish the two sources of uncertainty: its shape captures aleatory uncertainty, whereas its width captures epistemic uncertainty.

In this context, engineers are concerned with the contribution of the uncertainty in each input variable to the uncertainty in the QoI. Global sensitivity analysis provides a powerful set of tools to quantify the relative importance of each input parameter of the computational model with respect to the QoI. In particular, Sobol' indices are a popular variance-based global sensitivity measure, based on decomposing the variance of the output variable in terms of the variances of the input parameters (Sobol', 1993). However, Sobol' indices have been applied mainly to probabilistic variables rather than p-boxes.

In this paper, we extend the usage of Sobol' indices to the context of p-boxes. In particular, we make use of two recent developments. First, it has been shown in the recent literature, that the Sobol' indices can be estimated efficiently using sparse polynomial chaos expansions (PCE) (Sudret, 2008; Blatman and Sudret, 2010). PCE is a meta-modelling technique which approximates the computational model by a sum of weighted multivariate polynomials orthogonal with respect to the distributions of the input parameters (Blatman and Sudret, 2011). Second, Schöbi and Sudret (2015) introduced a generalized form of PCE suitable for p-boxes, based

on the definition of a suitable augmented space. In this contribution, we propose to build Sobol' indices for p-boxes by post-processing the corresponding augmented-space PCE.

The resulting Sobol' indices are interval-valued due to the definition of the input p-boxes. The interval width of each Sobol' index accounts for the epistemic uncertainty of the sensitivity measure. Hence, the interval width is in principle reducible by better characterizing the input variable. The intervals are an additional piece of information compared to the Sobol' indices based on probabilistic random variables. This information can be used for scheduling additional data generation for input variables when a constraint budget for additional measurements is available. Variables with large uncertainty bounds are preferred candidates for data enrichment.

The proposed approach is illustrated on a number of benchmark application examples. They show that an efficient estimation of PCE-based Sobol' indices is feasible in the context of p-boxes. Despite the increased complexity of the analysis, a small number of evaluations of the computational model is sufficient to estimate accurately the Sobol' indices.

References

- Blatman, G. and B. Sudret (2010). Efficient computation of global sensitivity indices using sparse polynomial chaos expansions. *Reliab. Eng. Sys. Safety* 95, 1216–1229.
- Blatman, G. and B. Sudret (2011). Adaptive sparse polynomial chaos expansion based on Least Angle Regression. *J. Comput. Phys* 230(6), 2345–2367.
- Ferson, S. and L. R. Ginzburg (1996). Different methods are needed to propagate ignorance and variability. *Reliab. Eng. Sys. Safety* 54(2-3), 133–144.
- Oberguggenberger, M., J. King, and B. Schmelzer (2009). Classical and imprecise probability methods for sensitivity analysis in engineering: A case study. *Int. J. Approx. Reason.* 50(4), 680–693.
- Schöbi, R. and B. Sudret (2015). Propagation of uncertainties modelled by parametric p-boxes using sparse polynomial chaos expansions. In *Proc. 12th Int. Conf. on Applications of Stat. and Prob. in Civil Engineering (ICASP12)*, Vancouver, Canada.
- Sobol', I. (1993). Sensitivity estimates for nonlinear mathematical models. *Math. Modeling & Comp. Exp.* 1, 407–414.
- Sudret, B. (2008). Global sensitivity analysis using polynomial chaos expansions. *Reliab. Eng. Sys. Safety* 93, 964–979.

A Bayesian based algorithm to build up sparse polynomial chaos expansions for global sensitivity analysis

Qian Shao^{1,2}, Thierry Mara^{*2}, and Anis Younes^{1,3}

¹Laboratoire Hydrologie et Geochemie de Strasbourg, University of Strasbourg/EOST, CNRS, 1 rue Blessig
67084 Strasbourg, France

²PIMENT, Université de La Réunion, 15 Avenue René Cassin, BP 7151, 97715 Moufia, La Réunion, France

³IRD UMR LISAH, F-92761 Montpellier, France

Global Sensitivity Analysis (GSA) aims at quantifying a mathematical model's output uncertainty due to changes of the input variables over the entire domain. Different methods have been developed in the last few decades to perform GSA [1], such as regression-based methods, variance-based methods, Morris method, sampling methods and so on. Among all these techniques, Sobol' sensitivity indices based on the ANalysis Of VAriance (ANOVA) [2] are of great interest by many researchers. These indices are usually computed by crude Monte Carlo simulation, which is computationally expensive and hardly applicable for complicated industrial models. To circumvent this problem, a metamodel with less expensive evaluations is usually adopted to substitute the original model under consideration for the computation of GSA.

In this context, the polynomial chaos expansion (PCE) using orthogonal polynomial bases has received much interest [3], with which the Sobol' indices can be computed exactly from algebraic operations on the coefficients of PCE. The computation of the PCE coefficients are often conducted in two approaches, the projection approach and the regression approach. The latter reveals efficient when dealing with a moderate number of input variables. However, with a large number of input variables or a high order of PCE, the number of coefficients increases dramatically, which requires a large number of model evaluations accordingly. To reduce the computational cost on direct model evaluations, one needs to decrease the number of coefficients in the PCE. To this aim, several approaches have been developed to construct a sparse PCE, where only basis functions and coefficients that have significant contributions to the variance of the model are retained. The original idea of sparse PCE came from Blatman and Sudret [4, 5], where they developed an iterative forward-backward algorithm to construct sparse PCE based on stepwise regression technique. Later, a least angle regression algorithm to build up sparse PCE was proposed by [6]. Hu and Youn [7] presented a sparse iterative scheme using the projection technique. More recently, Fajraoui et al. [8] developed a simple strategy to construct a sparse PCE using a fixed experimental design.

In this work, a new algorithm based on Bayesian Model Averaging (BMA) is proposed to construct the sparse PCE. The BMA, relying on Bayes' theorem, is a well-known statistical approach to

*Corresponding author: Thierry Mara, PIMENT, Université de La Réunion, 15 Avenue René Cassin, BP 7151, 97715 Moufia, La Réunion, France.

E-mail: thierry.mara@univ-reunion.fr

perform quantitative comparison of the proposed alternative models. The difficulty of the BMA lies in the evaluation of a quantity named Bayesian model evidence (BME), which involves an integral over the whole parameter space of a model, thus generally has no analytical solution. To approximate BME mathematically, the integral is treated with a Taylor series expansion followed by a Laplace approximation. Based on this approximation, the Kashyap information criterion (KIC) was proposed to evaluate BME for the most likely parameter set instead of the entire parameter space, making the evaluation computationally feasible. In the proposed algorithm, the maximum a posteriori estimate (MAP) is considered as the most likely coefficients for the selected model and used to evaluate KIC. For the sampling technique, the QMC method is adopted due to its space filling and desirable convergence properties. The proposed algorithm is outlined in the following:

- (i) An initial experimental design is generated with LPtau samples, followed by the model evaluations at the design points. Then a "full" PCE is constructed at given degree and interaction order.
- (ii) The basis functions in "full" PCE are reordered according to the contribution of each term to the variance of the model. This is performed in two steps: first, basis functions are sorted based on the Pearson correlation coefficients between basis functions and model evaluations, then they are reordered again by taking into account the partial correlation coefficients between basis functions and model evaluations.
- (iii) The sparse PCE is enriched by adding candidate basis polynomials with decreasing partial variance one by one. One eventually retains those terms that lead to a decrease of KIC to obtain an optimal sparse PCE model with minimized KIC.

The accuracy and efficiency of the proposed algorithm to construct sparse PCE for GSA is assessed by some benchmarking tests. The approach is then applied to perform the GSA on hydraulic models, which highlights the effectiveness and reliability of the proposed algorithm.

Bibliography

- [1] A. Saltelli, K. Chan, E. M. Scott, *et al.*, *Sensitivity analysis*, vol. 1. Wiley New York, 2000.
- [2] I. M. Sobol', On sensitivity estimation for nonlinear mathematical models, *Matematicheskoe Modelirovanie*, vol. 2, no. 1, pp. 112–118, 1990.
- [3] B. Sudret, Global sensitivity analysis using polynomial chaos expansions, *Reliability Engineering & System Safety*, vol. 93, no. 7, pp. 964–979, 2008.
- [4] G. Blatman and B. Sudret, Sparse polynomial chaos expansions and adaptive stochastic finite elements using a regression approach, *Comptes Rendus Mécanique*, vol. 336, no. 6, pp. 518–523, 2008.
- [5] G. Blatman and B. Sudret, An adaptive algorithm to build up sparse polynomial chaos expansions for stochastic finite element analysis, *Probabilistic Engineering Mechanics*, vol. 25, no. 2, pp. 183–197, 2010.
- [6] G. Blatman and B. Sudret, Adaptive sparse polynomial chaos expansion based on least angle regression, *Journal of Computational Physics*, vol. 230, no. 6, pp. 2345–2367, 2011.
- [7] C. Hu and B. D. Youn, Adaptive-sparse polynomial chaos expansion for reliability analysis and design of complex engineering systems, *Structural and Multidisciplinary Optimization*, vol. 43, p. 419–442, Sep 2010.
- [8] N. Fajraoui, T. A. Mara, A. Younes, and R. Bouhlila, Reactive transport parameter estimation and global sensitivity analysis using sparse polynomial chaos expansion, *Water Air Soil Pollut*, vol. 223, p. 4183–4197, May 2012.

Identification of influential parameters in building energy simulation and life cycle assessment

S. M. Spiessl¹ and D.-A. Becker¹

¹Gesellschaft fuer Anlagen - und Reaktorsicherheit (GRS) gGmbH, Germany

Abstract

To support the safety assessment of final repositories for radioactive waste, it is very helpful to do sensitivity analysis (SA). The applied numerical models, however, may exhibit a highly non-linear behaviour. For example, quasi-discontinuous behaviour may occur when barriers fail to function at some point in time, which can happen very fast. This can cause a two-split output distribution. In addition, the model output is dominantly very low and varies over many orders of magnitude, which may result in a steep and asymmetric distribution of the model output. Such distributions and model behaviours are challenges for performing proper SA.

All methods of SA have their specific advantages and disadvantages in view of the nature of the model under consideration, number of parameters, input distributions and available computational power. Furthermore, SA will also depend upon which questions need to be answered about the sensitivity of the system. Objectives of a guideline for SA are that the analyses produce unique and robust results as well as provide clear answers to the asked questions. The guideline presented in this talk is a recommendation based on the outcome of detailed experiments with different methods and sample sizes.

We generally recommend using a quasi-Monte Carlo sampling method for SA as it covers the parameter space more homogeneously and produces more robust results compared to random sampling and existing samples can be extended. Although, from a radiological point of view, it may be sufficient to only analyse the peak values of all runs, we recommend performing a time-dependent analysis in addition, as it provides more insight into the system behaviour.

In general, the nature of the numerical model will not be fully known before the analysis. Therefore, a stepwise approach is advisable, starting with graphical SA and then advancing to more sophisticated methods. This approach may also allow detecting sensitivities which are not found with one type of method.

Graphical methods provide a visual insight into the model behaviour and sensitivity of the different parameters. They can be used as the first approach to identify potentially important parameters. Graphical SA is in particular useful when many parameters need to be considered. The number of parameters may be reduced in this first analysis which makes more advanced methods easier to use and less time-consuming in terms of required computational power. With the graphical methods Contribution to the Sample Mean (CSM) plot and scatterplots we obtained good results.

In the next step, regression- or correlation-based methods are suggested to be applied. This can be done on a value- or rank-basis. If nonlinear effects play a role, the rank-based version is often more adequate, but that has to be checked as the coefficient of determination (R^2) can even decrease under a rank transformation. We made good experience with the standard regression method SRC and standard rank regression method SRRC.

Variance-based methods of SA do not require linearity or monotonicity, which does, however, not automatically mean that they work better on a nonlinear model. As a first approach of this kind we recommend applying the EASI algorithm, which is numerically effective, can be used with any sample and seems to yield robust results. However, EASI can only compute the first-order effects. A low sum of all first-order sensitivity indices indicates poor significance of the first-order analysis. This can often be improved by applying an adequate transformation to the model output. Otherwise, an analysis of the higher orders is recommended in such cases.

If certain model properties are known or have been identified during the above analysis, it is recommended to add specifically designed investigations. For instance, influences of parameter correlations may be seen in the difference between the Partial Correlation Coefficients (PCC) and SRC results or Partial Rank Correlation Coefficients (PRCC) and SRRC results. Non-parametric methods like the two-sample Smirnov test can also be helpful in specific cases.

Acknowledgements: This work was funded by the German Federal Ministry for Economic Affairs and Energy (BMWi) under grant No. 02E10941.

Perturbed-Law based sensitivity Indices for sensitivity analysis in structural reliability

ROMAN, SUEUR[†]; NICOLAS, BOUSQUET[†]; BERTRAND, IOOSS[†];
JULIEN, BECT[‡]

[†]EDF Lab Chatou, France; [‡]CentraleSupélec, France

For sensitivity analysis of model outputs, the most popular methods are those based on the variance decomposition of the output, such as the Sobol' indices, as they allow defining easy-to-interpret indices measuring the contribution of each input variable to the overall output dispersion. However, these indices are not adapted to the analysis of the impact of inputs on quantities characterizing extreme events, such as a failure probability, a quantile or a failure domain (Lemaître et al., 2015). We denote $g(\cdot)$ the studied model, $\mathbf{X} = (X_1, \dots, X_d) \in \mathbb{X}$ the random vector of the d independant input variables and $Y = g(\mathbf{X})$ the model output. $\mathbf{X}$ follows a joint probability density function $f(\mathbf{x})$. We focus on a typical problem in structural reliability (Morio and Balesdent, 2015), which is the analysis of the failure probability (the failure event arrives with the event $g(\mathbf{x}) < 0$):

$$p = \int_{\mathbb{X}} \mathbf{1}_{\{g(\mathbf{x}) < 0\}} f(\mathbf{x}) d\mathbf{x}.$$

PERTURBED-LAW BASED SENSITIVITY INDICES

Based on the perturbation of each marginal input density, a new type of sensitivity indices has been recently developed (Lemaître et al., 2015). An input X_i , with marginal density f_i , is replaced by the perturbed variable $X_{i\delta}$. $X_{i\delta}$ follows the density $f_{i\delta}$, based on the initial f_i perturbed of a δ quantity. This allows defining the following perturbed probability $p_{i\delta}$:

$$p_{i\delta} = \int_{\mathbb{X}} \mathbf{1}_{\{g(\mathbf{x}) < 0\}} \frac{f_{i\delta}(x_i)}{f_i(x_i)} f(\mathbf{x}) d\mathbf{x}.$$

From this quantity, we define the Perturbed-Law based Indices (PLI) in the following way:

$$S_{i\delta} = \left[\frac{p_{i\delta}}{p} - 1 \right] \mathbf{1}_{\{p_{i\delta} > p\}} + \left[\frac{p}{p_{i\delta}} - 1 \right] \mathbf{1}_{\{p_{i\delta} \leq p\}}. \quad (1)$$

The PLI measures have some expected properties, such as being equal to 0 when the failure probability is not changed by the perturbation, or taking a sign that indicates the direction of change of the probability with the δ perturbation. $f_{i\delta}$ is obtained by minimizing the Kullback-Leibler divergence KL between $f_{i\delta}$ and f_i for a given shift δ of a statistical characteristic (for example the mean, the variance, a quantile, ...) of the distribution of X_i :

$$KL(f_{i,\delta}, f_i) = \int_{-\infty}^{+\infty} f_{i,\delta}(x_i) \log \frac{f_{i,\delta}(x_i)}{f_i(x_i)} dx_i.$$

Easy minimization of KL provides explicit (analytical) solutions of $f_{i\delta}$ for a large range of perturbed parameters of f_i (e.g. mean, variance, ...) on classical pdf (e.g. Gaussian). In other cases, a numerical optimization procedure is required (Lemaître et al., 2015).

MONTE-CARLO ESTIMATION OF THE PLI MEASURES

An analytical calculation of the probabilities p and $p_{i\delta}$ is not possible in practice, and a Monte-Carlo sample $(\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(N)})$ is used. Thus, the estimations are respectively:

$$\hat{p}_N = \frac{1}{N} \sum_{n=1}^N \mathbf{1}_{g(\mathbf{x}^n) \leq 0} \text{ and } \hat{p}_{i\delta, N} = \frac{1}{N} \sum_{n=1}^N \mathbf{1}_{\{g(\mathbf{x}^n) \leq 0\}} \frac{f_{i\delta}(x_i^n)}{f_i(x_i^n)}.$$

In practice, the estimation of PLI might require a large-size sample if the failure probability is very low (for instance 10^{-6} or less). This could reveal impractical if the code g is costly.

To improve the estimation of p and $p_{i\delta}$, the importance sampling or subset simulation methods could be used but they often remain too costly in practice (Morio and Balesdent, 2015). We propose here to use a metamodel-based approach which has shown a noticeable efficiency for Sobol' indices (Le Gratiet et al., 2017). In particular, the Gaussian process metamodel allows to control the error on the sensitivity estimates due to the metamodel approximation.

BAYESIAN IMPORTANCE SAMPLING

The so-called Bayesian Importance Sampling (BIS) consists in two steps :

- 1) A Gaussian process model $\tilde{g}$ of the code is built using a budget of $N_0 < N$ computer experiments. This allows to define a relevant importance density at the following step;
- 2) p and $p_{i\delta}$ are estimated with a $N_1 = N - N_0$ importance sampling scheme (Bect et al., 2015).

The Bayesian optimal importance density $f_{IS}^*(\mathbf{x})$ is proportional to $f(\mathbf{x})\sqrt{\mathbb{P}_{N_0}^{\tilde{g}}[g(\mathbf{x}) < 0]}$. If $\hat{Z}$ is an estimator of $Z = \int_{\mathbb{X}} f(\mathbf{x})\sqrt{\mathbb{P}_{N_0}^{\tilde{g}}[g(\mathbf{x}) < 0]}d\mathbf{x}$, the estimate of p is

$$\hat{p}_{i\delta, N_0, N_1}^{BIS} = \frac{\hat{Z}}{N_1} \sum_{n=1}^{N_1} \frac{\mathbf{1}_{\{\mathbb{E}_{N_0}^{\tilde{g}}[\tilde{g}(\mathbf{x}^n)] < 0\}}}{\sqrt{\mathbb{P}_{N_0}^{\tilde{g}}[g(\mathbf{x}^n) < 0]}} \frac{f_{i\delta}(x_i^n)}{f_i(x_i^n)}.$$

PERSPECTIVES: PLI FOR SENSITIVITY ANALYSIS OVER A QUANTILE

In safety studies, a large quantile value is often preferred to a failure probability computation. PLI measures can be applied to an output quantile $q_\alpha(Y) = \inf\{y \text{ s.t. } F(y) \geq \alpha\}$ with F the distribution function of Y . This requires to combine quantile estimation and importance sampling. A naive approach for the estimation of quantiles could consist in replacing F by its empirical estimator in the latest formula. By re-ordering the $y^n = g(\mathbf{x}^n)$ we could define $k_{\alpha, i\delta}$ by

$$k_{\alpha, i\delta} = \min \left\{ k \text{ s.t. } \left[\frac{\sum_{n=1}^k f_{i\delta}(x_i^{(n)})}{f_i(x_i^{(n)})} \right] / \left[\frac{\sum_{n=1}^N f_{i\delta}(x_i^{(n)})}{f_i(x_i^{(n)})} \right] \geq \alpha \right\},$$

where each (n) denotes the re-ordered index of y^n in our sample. A straightforward estimator of $q_\alpha(Y)$ is given by $\hat{q}_\alpha(Y) = y^{k_\alpha}$. However, this estimator is unlikely to show good consistency properties, since it is built on a non-continuous quantile function obtained by inverting the empirical cumulative distribution function of Y . A more promising approach would be to use a quantile-regression framework (Egloff and Leippold, 2010). We notice that $q_\alpha(Y)$ can be written as $\operatorname{argmin}_r \mathbb{E}_X[\rho_\alpha(g(X) - r)]$, where ρ_α denotes the check-function $\rho_\alpha(u) = u(\alpha - \mathbf{1}_{\{u < 0\}})$ (see Fort et al., 2016). This suggests the following estimator:

$$\hat{q}_{\alpha, i\delta}(Y) = \operatorname{argmin}_r \sum_{n=1}^N \rho_\alpha(g(\mathbf{x}^n) - r) \frac{f_{i\delta}(x_i^n)}{f_i(x_i^n)}.$$

Replacing p and $p_{i\delta}$ in (1) respectively by $\hat{q}_\alpha$ and $\hat{q}_{\alpha, i\delta}$ provides the wanted PLI over a quantile.

References:

- J. Bect, R. Sueur, A. Gérossier, L. Mongellaz, S. Petit and E. Vazquez (2015), *Échantillonnage préférentiel et méta-modèles : méthodes bayésiennes optimale et défensive*, 47èmes Journées de Statistique de la SFdS, Lille, FR.
- D. Egloff, M. Leippold (2010), *Quantile estimation with adaptive importance sampling*, Ann. Statist., 38:1244-1278.
- L. Le Gratiet, S. Marelli and B. Sudret (2017), Metamodel-Based Sensitivity Analysis: Polynomial Chaos Expansions and Gaussian Processes, In: *Springer Handbook on UQ*, R. Ghanem, D. Higdon and H. Owhadi (Eds).
- P. Lemaître, E. Sergienko, A. Arnaud, N. Bousquet, F. Gamboa and B. Iooss (2015), Density modification based reliability sensitivity analysis, *Journal of Statistical Computation and Simulation*, 85:1200-1223.
- J. Morio et M. Balesdent (2015), *Estimation of Rare Event Probabilities in Complex Aerospace and Other Systems, A Practical Approach*, Woodhead Publishing.
- J. Fort, T. Klein, N. Rachdi (2016), *New sensitivity analysis subordinated to a contrast* Communication in Statistics : Theory and Methods. 45:4349-4364.

[Roman Sueur; EDF R&D, 6 Quai Watier, 78401 Chatou, France]
 [roman.sueur@edf.fr -]

VARIANCE-BASED SENSITIVITY INDICES FOR INPUTS DEFINED OVER NON-RECTANGULAR DOMAINS.

Stefano Tarantola, Biagio Ciuffo and Qiao Ge

Most of the literature on variance-based methods for sensitivity analysis focuses on independent inputs. In recent years, interest has increased for correlated input. However, in both cases the domain of the inputs is rectangular, such as a hypercube of size one. In practical cases, it happens that the domain of the inputs is not rectangular but has a different shape, such as a circular area, or a triangle.

Let us denote by $x = (y, z)$ a vector of inputs of dimension d where y and z are two sub-sets of inputs of size $s < d$ and $d-s$. x is assumed distributed with probability density function $p(y, z)$ over the unit hypercube $I^d(0,1)$. The marginal distribution of y is indicated with $p(y)$ and the conditional distribution of y given z is indicated with $p(y|z)$. The output is assumed scalar and is given by $f(y, z)$, meaning that it is obtained by evaluating the function f at point x .

First order S_y and total order S_y^T sensitivity indices for input y are defined in [1]:

$$S_y = \frac{1}{D} \left[\int_{R^n} f(y', z') p(y', z') dy' dz' \left[\int_{R^{n-s}} f(y', \hat{z}) p(y', \hat{z} | y') d\hat{z} - \int_{R^n} f(y, z) p(y, z) dy dz \right] \right]$$

$$S_y^T = \frac{1}{2D} \int_{R^{n+s}} [f(y, z) - f(\bar{y}', z)]^2 p(y, z) p(\bar{y}', z | z) dy d\bar{y}' dz$$

where (y, z) and (y', z') are two independent random vectors generated from the joint distribution $p(y, z)$, $\hat{z}$ is a random vector generated from the conditional probability density function $p(y', \hat{z} | y')$ and $(\bar{y}', z)$ is a random vector generated from the conditional probability density function $p(\bar{y}', z | z)$.

In the present study we propose four approaches to compute variance-based sensitivity indices in the case of non-rectangular domains and compare their performance. As benchmark, we consider the g -function in two settings and we analyse two types of non-rectangular domain: circular, at different radii, and triangular in different conditions. The analytic sensitivity indices for all case studies have been computed using the formulas above and have been used as the reference for the benchmark.

- The first approach uses the classic Sobol' formula for the uncorrelated case as proposed in [1] coupled with a rejection method that considers the points only if they are inside the domain. In this case, we test to what extent this formula is valid as the domain gets less and less rectangular.
- The second approach uses the Sobol' formula for correlated inputs proposed by [2] coupled with a rejection method proposed in [3].
- The third and fourth approaches evaluate the sensitivity indices using the definition, i.e. computing the variance of the conditional expectation of the model output, given one of the inputs, and dividing by the total output variance. Both approaches start from a set of Monte Carlo points and divide the range of a given input in bins.

While in the third method the same number of points in each bin is considered (implying that the bins can have different size), the fourth method uses the same bin size.

References

- [1] Saltelli, A., P. Annoni, I. Azzini, F. Campolongo, M. Ratto, S. Tarantola (2010) Variance based sensitivity analysis of model output. Design and estimator for the total sensitivity index, *Computer Physics Communications*, 181, 259–270
- [2] Kucherenko S, S. Tarantola and P. Annoni(2012) Estimation of global sensitivity indices for models with dependent variables, *Computer Physics Communications*, Vol: 183, Pages: 937-946
- [3] Kucherenko, S., O.V. Klymenko and N. Shah (2016) Sobol' indices for problems defined in non-rectangular domains, Submitted to *Comp. Phys. Comm.*

Confidence intervals for Sobol' indices

TAIEB, TOUATI

Pierre and Marie Curie University, France

When studying the interactions between variables, going beyond regressions is crucial for more precision and accuracy. In fact, studying the interdependence between the variances of each of the model's components adds a wider array of interpretation and forecasting techniques. The literature defines this analysis segment as variance based sensitivity analysis which is regarded as one of the most frequently used computer models in engineering studies (Ferretti et al., 2016). The model's output variance that is caused by a specific model input or a combination of more than one input (Sobol, 1993; Iooss et al., 2015).

$$S_i = \frac{V_i}{V} = \frac{\text{Var}[\mathbb{E}(f(X)|X_i)]}{\text{Var}[f(X)]} \text{ and } S_i^{\text{tot}} = \frac{V_i^{\text{tot}}}{V} = 1 - \frac{V_{-i}}{V} = 1 - \frac{\text{Var}[\mathbb{E}(f(X)|X_{-i})]}{\text{Var}[f(X)]}, \quad (1)$$

where $f(X)$ is the computer model, $X = (X_1, \dots, X_d) \in \mathbb{R}^d$ are the model inputs (independent random variables), $i = 1, \dots, d$, and X_{-i} is the input vector except X_i . S_i , the first-order Sobol' index, only includes the sole effect of X_i , while S_i^{tot} , the total Sobol' index, takes into account all the effects of X_i including its interaction effects with other inputs. For u a subset of $\{1, 2, \dots, d\}$ we consider the partition: $X = X_u \cup X_{\bar{u}}$, where $\bar{u}$ is the complement of u in $\{1, 2, \dots, d\}$.

As in Iooss et al. (2016), we chose to study estimators which provide $(\hat{S}_i, \hat{S}_i^{\text{tot}})$, estimates of (S_i, S_i^{tot}) , by using two independent input designs **A** and **B**, matrices with n rows (sample size) and d columns. We focus especially on the Martinez estimator that sets Sobol indices as correlation coefficient. The mathematical properties of the empirical correlation coefficient lead to explore more thoroughly the properties of this estimator and build thereafter confidence intervals. Martinez estimator (Martinez, 2011): By noticing that

$$S_i = \rho(f(\mathbf{B}), f(\mathbf{A}_{B(i)})) \text{ and } S_i^{\text{tot}} = 1 - \rho(f(\mathbf{A}), f(\mathbf{A}_{B(i)})) \quad (2)$$

where $A_{B(u)} = A_u \cup B_{\bar{u}}$, u a subset of $\{1, 2, \dots, d\}$ (for Martinez estimator $u=i$, $i = 1, \dots, d$). ρ is the linear correlation coefficient, the Sobol' indices can be estimated using the well-conditioned empirical formula of ρ (i.e. using the product of differences).

For the Martinez estimator, asymptotic confidence intervals are approximated by using Fisher's transformation applied to the sample correlation coefficients $\hat{S}_i$ and $\hat{S}_i^{\text{tot}}$ from Eq 2. It is only valid under Gaussian hypothesis of the output variable distribution. The classical 95% confidence intervals obtained by the Martinez method are described in Iooss et al. (2016).

Based on the fact that the Sobol indices are interpreted as correlation coefficients, we give two asymptotic results which will be applied for Sobol indices. This methodology is analogue to the demonstration given by Lehman (1999). We provide a formula for the asymptotic variance as a polynomial function of the correlation coefficient.

We assume that (Y, Z) is a squared integrable couple of random variables. R_n is the empirical correlation coefficient of (Y, Z) and ρ the theoretical correlation coefficient. $C_n, \sigma_n(Y)$ and $\sigma_n(Z)$ mean respectively the empirical covariance and the empirical variances. The first theorem concerns the asymptotic normality of the triplet $\{C_n, \sigma_n(Y), \sigma_n(Z)\}$. If there K is the covariance matrix formed after applying the central limit theorem to the triplet $\{C_n, \sigma_n(Y), \sigma_n(Z)\}$. the asymptotic normality of R_n gives:

$$\sqrt{n}(R_n - \rho) \rightarrow \mathcal{N}(0, \tau^2) \quad (3)$$

τ^2 is a polynomial function of ρ , this can facilitate the implementation of the method. $\tau^2 = P(\rho)$ where:

$$P(x) = Ax^2 + Bx + C \quad (4)$$

A, B and C depends on the coefficients of K .

Remark

Bishara et al. (2016) gives recently several alternatives to Fisher's method to compute confidence intervals when data are not normal. The methods are classified in two main groups: Transforming data and Bootstrapping. For the transforming data methods the best performance was performed by the well known Speraman rank-order and the rank inverse normal transformation. Among the bootstrapping methods Efron et al. (1994), which have the merit of conserving the original scale of raw data, an observed imposed bootstrap had an adequate coverage probability with precise intervals comparing to other Bootstrap methods.

The work of Beasley et al. (2007) served as a foundation for this method in which computing time has been reduced making computations easier for larger samples.

In this communication, the extension of the Martinez method to non Gaussian distribution is studied. Indeed, non Gaussianity can distort the Fisher's confidence interval, and the outcome can be quite misleading. The two following points will be discussed:

1. Asymptotic confidence intervals. In this case, through the methodology described in Remark 2 we give an asymptotic confidence interval for Sobol' indices in a general case.
2. Non asymptotic confidence intervals. In this case, we compare several methods to improve the Martinez method while keeping the approximation approach on the one hand and with a Bootstrapping approach on the other hand. We base this study on the methodology described in remark 3. Comparisons are made in terms of coverage probability and confidence interval length.

Numerical studies will illustrate all these effects for the different methods, demonstrating that with the asymptotic method we have more accurate coverage probability comparing to the Martinez approach. The results suggest that sample non Gaussianity can justify avoidance of the Fisher's confidence interval in favor of more robust alternative (Non-asymptotic or asymptotic).

References:

- F. Ferretti, A. Saltelli and S. Tarantola (2016), Trends in sensitivity analysis practice in the last decade, *Science of the Total Environment*, in press.
- JM. Martinez (2011), Analyse de sensibilité globale par décomposition de la variance, *Presentation in "Journée des GdR Ondes & Mascot Num"*, 13 janvier 2011, Institut Henri Poincaré.
- B.Iooss and P.Lemaître (2015), A review on global sensitivity analysis methods. *Uncertainty Management in Simulation-Optimization of Complex Systems*. 101–122, Springer.
- I. Sobol (1993), Sensitivity estimates for non linear mathematical models. *Mathematical Modelling and Computational Experiments*, 1:407-414.
- E.L Lehman(1999), Elements of large-sample theory. *Springer Science & Business Media*.
- A J.Bishara and J B.Hittner (2016), Confidence intervals for correlations when data are not normal. *Behavior Research Methods*, in press.
- B.Efron and RJ.Tibshirani (1994), An introduction to the bootstrap. *CRC press*
- WH.Beasley, L.DeShea, LE.Toothaker, JL.Mendoza, DE.Bard, and JL.Rodgers (2007), Bootstrapping to test for nonzero population correlation coefficients using univariate sampling. *Psychological methods American Psychological Association.*, 12:414.
- B. Iooss, M.Baudin, K.Boumhaout, T.Delage, J.M Martinez (2016), Numerical stability of Sobolà indices estimation formula. *Submitted to SAMO 2016 Conference.*, La Réunion, France.
- [Taieb Touati; UPMC, 4 Place Jussieu, 75005 Paris, France]
[taieb.touati@etu.upmc.fr – <https://cran.r-project.org/web/packages/sensitivity/index.html>]

Global sensitivity analysis by HDMR combining with the improved GMDH algorithm

Lu Wang^a, Shufang Song^b, Shiyu Liu^c

(School of Aeronautics, Northwestern Polytechnical University, Xi'an, Shaanxi, 710072, PR China)

^alouisewanglu@gmail.com ^bshufangsong@nwpu.edu.cn ^cshiyu_liu22@163.com

Abstract

Global sensitivity analysis (GSA) is a very useful tool to evaluate the influence of input variables in the whole distribution range. The interactive influences have been taken into consideration in GSA. The variance-based GSA was proposed by Sobol' according to ANOVA decomposition^[1]:

$$f(\mathbf{x}) = f_0 + \sum_{i=1}^n f_i(x_i) + \sum_{1 \leq i < j \leq n} f_{ij}(x_i, x_j) + \sum_{1 \leq i < j < k \leq n} f_{ijk}(x_i, x_j, x_k) + \dots \quad (1)$$

where the input variables are $\mathbf{x}=(x_1, x_2, \dots, x_n)^T$. Furthermore, through a complete basis set of orthonormal polynomials function, random sampling high dimensional model representation (HDMR) can be used to solve the Sobol' first and second order global sensitivity indices^[2].

$$f_i(x_i) \approx \sum_{p=1}^{m_1} \alpha_p^i \varphi_p(x_i), \quad f_{ij}(x_i, x_j) \approx \sum_{p=1}^{m_2} \sum_{q=1}^{m_3} \beta_{pq}^{ij} \varphi_{pq}(x_i, x_j) \quad (2)$$

where m_1, m_2, m_3 are suitable maximum orders. And coefficients $\alpha_p^i, \beta_{pq}^{ij}$ are calculated by random sampling as follows:

$$\alpha_p^i \approx \frac{1}{N} \sum_{s=1}^N f(\mathbf{x}^{(s)}) \varphi_p(x_i^{(s)}), \quad \beta_{pq}^{ij} \approx \frac{1}{N} \sum_{s=1}^N f(\mathbf{x}^{(s)}) \varphi_{pq}(x_i^{(s)}, x_j^{(s)}) \quad (3)$$

However, the flaws of classical HDMR method cannot be ignored. For instance, it needs a large number of samples N to calculate the decomposition coefficients and cannot calculate high order sensitivity indices. The group method of data handling(GMDH) is a family of inductive algorithms for computer-based mathematical modeling of multi-parametric datasets that features fully automatic structural and parametric optimization of models. In order to improve the GMDH, neural network algorithm are integrated to generate the assured function description of the nonlinear model with high precision using a relative small samples^[3]. Besides, the IGMDH algorithm has a fixed number of layers which cannot exceed 3.

Thus, we consider to combine the improved group method of data handling (IGMDH) with HDMR for calculating the coefficients more efficiently. The main procedures of the IGMDH-HDMR method that we proposed are presented in Figure 1, and the structure of it is shown in Figure 2. The metamodeling of IGMDH-HDMR is:

$$f(\mathbf{x}) = f_0 + \sum_p \alpha_p^i \varphi_p(x_i) + \sum_{p,q} \beta_{pq}^{ij} \varphi_{pq}(x_i, x_j) + \sum_{p,q,r} \gamma_{pqr}^{ijk} \varphi_{pqr}(x_i, x_j, x_k) + \sum_{p,q,r,s} \delta_{pqrs}^{ijkl} \varphi_{pqrs}(x_i, x_j, x_k, x_l) \quad (4)$$

As a result, Sobol' sensitivity indices can be calculated by the derived coefficients:

$$\hat{S}_i = \sum_p (\alpha_p^i)^2 / \hat{D}, \quad \hat{S}_{ij} = \sum_{p,q} (\beta_{pq}^{ij})^2 / \hat{D}, \quad \hat{S}_{ijk} = \sum_{p,q,r} (\gamma_{pqr}^{ijk})^2 / \hat{D}_n, \quad \hat{S}_{ijkl} = \sum_{p,q,r,s} (\delta_{pqrs}^{ijkl})^2 / \hat{D}_n$$

$$(5) \hat{S}_i^{tot} = \hat{S}_i + \sum_{j \neq i} \hat{S}_{ij} + \sum_{k \neq j \neq i} \hat{S}_{ijk} + \sum_{l \neq k \neq j \neq i} \hat{S}_{ijkl} \quad (6)$$

where total variance is calculated as:

$$\hat{D} = \sum_p (\alpha_p^i)^2 + \sum_{p,q} (\beta_{pq}^{ij})^2 + \sum_{p,q,r} (\gamma_{pqr}^{ijk})^2 + \sum_{p,q,r,s} (\delta_{pqrs}^{ijkl})^2 \quad (7)$$

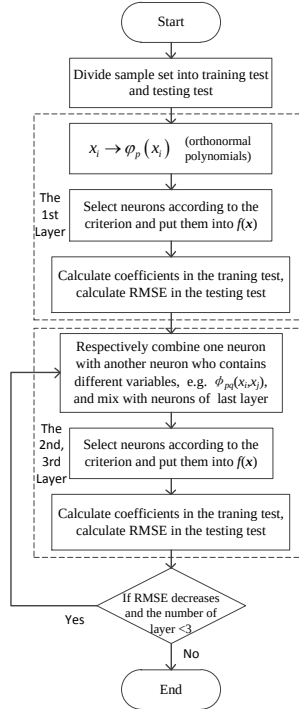


Fig. 1 The flowchart of GMDH-HDMR

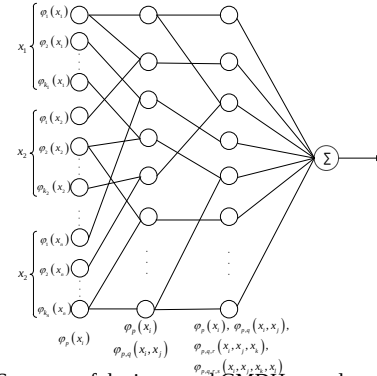


Fig. 2 Structure of the improved GMDH neural network

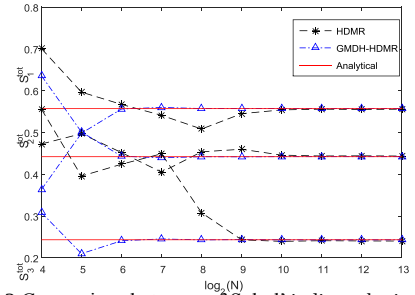


Fig. 3 Comparison between Sobol' indices obtained from GMDH-HDMR with RS-HDMR

Here are the details of how to choose the maximal order. Assume $p_{max} = m_1$ and calculate the global sensitivity at different order. Thus we can get a series of $\hat{S}_i^p, \hat{S}_{ij}^{pq}, \hat{S}_{ijk}^{pqr}, \hat{S}_{ijkl}^{pqrs}$.

When $|\hat{S}_i^{p^*+1} - \hat{S}_i^{p^*}| < \varepsilon_1, p_{max} = p^*$. When $|\hat{S}_{ij}^{p^*+1, q^*+1} - \hat{S}_{ij}^{p^*, q^*}| < \varepsilon_2, (p_{max}, q_{max}) = (p^*, q^*)$.

When $|\hat{S}_{ijk}^{p^*+1, q^*+1, r^*+1} - \hat{S}_{ijk}^{p^*, q^*, r^*}| < \varepsilon_3, (p_{max}, q_{max}, r_{max}) = (p^*, q^*, r^*)$.

When $|\hat{S}_{ijkl}^{p^*+1, q^*+1, r^*+1, s^*+1} - \hat{S}_{ijkl}^{p^*, q^*, r^*, s^*}| < \varepsilon_4, (p_{max}, q_{max}, r_{max}, s_{max}) = (p^*, q^*, r^*, s^*)$

The neurons whose order exceed the maximal value will be deleted.

Finally, the Ishigami function is considered, which is a highly nonlinear function of three inputs: $f(x) = \sin(x_1) + 7\sin^2(x_2) + 0.1x_3^4\sin(x_1)$, where $x_i (i=1,2,3)$ are uniformly distributed on the interval $[-\pi, \pi]$. And the comparisons of Sobol' indices obtained from IGMDH-HDMR with HDMR are listed in Figure 3. It is obvious that both precision and convergence of the IGMDH-HDMR method are higher than those of the RS-HDMR method.

Keywords: global sensitivity analysis; high dimensional model representation(HDMR); group method of data handling(GMDH) algorithm; Sobol' sensitivity indices

References

- [1] I. M. Sobol. Global sensitivity indices for nonlinear mathematical models and their Monte Carlo estimates. *Mathematics and Computers in Simulation*, 2001, 55(1): 271–280.
- [2] B. Feil, S. Kucherenko, N. Shah. Comparison of Monte Carlo and quasi Monte Carlo sampling methods in high dimensional model representation. *First International Conference on Advances in System Simulation*, Portugal, 2009: 12-17.
- [3] AG. Ivakhnenko. Heuristic self-organization in problems of engineering cybernetics. *Automatica*, 1970, 6(2): 207-219.

A new method of network clustering based on second-order sensitivity index

Huan Liu, Qiongli Wu*, Yiming Ding

Wuhan Institute of Physics and Mathematics, Chinese Academy of Sciences

*wuqiongli@wipm.ac.cn

Networks are organized into communities with dense internal connections, giving rise to high values of the clustering coefficient. In addition, these networks have been observed to be assortative, i.e., highly connected vertices tend to connect to other highly connected vertices, and have broad degree distributions. It means that high interaction between vertices exist within one cluster.

For such an important issue of interaction identification between factors, second-order Sobol's index provides a quantitative way to reveal the case. However, for decades after Sobol's method was proposed, the sensitivity analysis (SA) work using Sobol's indices mostly focuses on the aim of factor prioritization (FP) using first-order Sobol's index and factor fixing (FF) using total-order Sobol's index. The application is only qualitatively limited to check which pair of factors' interaction is 'larger' and which pair is 'smaller'. What's more, the performance after these interaction mapping has been ignored. In such case, the potential usefulness of identifying the interaction quantitatively for second-order Sobol's index is not fully tapped.

We have been clear that the interaction reflects the effect strength between factors. It is the intrinsic property of the system. By coincidence, the effect strength between factors is the basis of constructing a network. In a network, communities can be defined as sets of vertices with dense internal connections, such that the inter-community connections are relatively sparse. It is similar to the factor clustering in [1]. Factors are clustered into several modules according to expert modelling experience, then the interaction between modules are detected by Sobol's group analysis. The internal-module interaction is dense and the inter-module interaction is relatively sparse.

To understand the network behaviors, such as the small-world property and high degree of clustering, we need to detect community structure in networks. Network clustering is vastly used for such structure detection. Again, such network clustering needs the quantitative description of the effect strength between factors, and the task can be fulfilled by second-order sobol's index. It means that we can use the interaction information given by second-order sobol's index for network clustering, so that to detect model clustering structure. If such method works, we can even detect the module structure in the work of [1] without expert modelling experience, which problem is also listed in the discussion of the work. The significance of such method is that it can providing a configuration of putting factors into naturally clustering according to the intrinsic interaction property of the model itself, not the subjective knowledge of the modellers. The uncertainty of expert knowledge could be avoid in this way, because if the clustering work by expert knowledge is not validated, the modules analysis result could be misleading.

As such, we propose a new method of network clustering based on second-order sensitivity index in this paper. This method is a combination of second-order Sobol's index matrix and network clustering method.

We use Newman's fast algorithm for detecting community structure in networks. Newman's fast algorithm is an agglomerative algorithm, and its basic idea is from the greedy algorithm. It firstly initializes a network with n vertices, and assumes there are m edges in the network. The fraction of edges e_{ij} is decided by the second-order sensitivity indices of the factors. Then we join communities together in pairs if they have edges connected. At the same time calculate the increase of the modularity Q :

$$\Delta Q = e_{ij} + e_{ji} - 2 \times a_i \times a_j = 2(e_{ij} - a_i \times a_j)$$

Choosing at each step the join that results in the greatest increase (or smallest decrease) in Q . In the process of joining, update the corresponding element e_{ij} , and add the row and columns associated to i, j in the same community in order. Then repeat the step, until the whole network join into a bigger community, then stop.

We present here a test case on a specific function constructed on the basis of G-function:

$$F = g(x_1, x_2, \dots, x_5) + g(x_6, x_7, \dots, x_{10}) + g_1(x_1) \cdot g_6(x_6) \cdot g_7(x_7) + g_1(x_1) \cdot g_2(x_2) \cdot g_6(x_6)$$

We can see the main part of F are $g(x_1, x_2, \dots, x_5)$ and $g(x_6, x_7, \dots, x_{10})$, and there are also two items constructed by the multiplication of factors from these two items, namely there is connection between the two communities in the network structure, which makes the network clustering not so obvious.

We computed the second-order indices of all the factors, and constructed its network structure as shown in Fig.1(a). Network clustering analysis is shown in Fig.2(b). We can see two groups were separated obviously by our method, which are $(x_1, x_2, \dots, x_5)$ and $(x_6, x_7, \dots, x_{10})$. It means that, even with some connections introduced by the multiplication of some vertices, our method can still make a good network clustering.

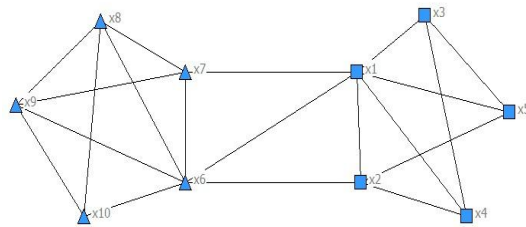


Fig.1.(a)

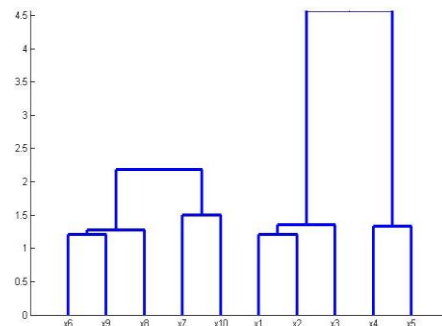


Fig.1.(b)

Our method can also make factor clustering 'blindly' to expert knowledge of modelling while revealing the model mechanisms in the case that sensitivity analysis needs to be performed to group of factors. An application to an ecological model will be given in the full paper to present the factor clustering by our method in model analysis processing.

References:

- [1] Wu, Q. and Cournède, P.-H., A comprehensive methodology of global sensitivity analysis for complex mechanistic models with an application to plant growth. Ecological complexity, 20:219-232.

Model Interpretation via Multivariate Padé-Approximant in Low-Rank Format

Rastko Živanović (rastko@eleceng.adelaide.edu.au and zivanovr@yahoo.com)
The University of Adelaide, Australia

Summary:

The aim of this paper is to present a computationally efficient method for approximation of a response function in several variables using a novel class of multivariate rational functions in low-rank format, constructed via power series summation. This method is well suited for high-dimensional problems where evaluating response function will take long time. In addition to sensitivity analysis, models based on the rational functions, are able to provide further insight into response of a man-made systems (or physical phenomena) under investigation; e.g. it is possible to reconstruct response hypersurfaces beyond region that was used to build the approximation, and certain types of singularities can be located with high accuracy.

The concept of Padé-approximants, which is the kernel of the proposed method, has been already generalized to multivariate function approximation [1]. It is well known that such approximants, applied in summation of power series, are useful for interpretation of models (i.e. parametrized response functions). “Magically”, it is possible to deduce important properties of a model from a given power series. For example, local behavior near a singularity, characterized with parameter values for which a model is unbounded or not unique, can be analyzed. It is important to note that Padé approximation is the optimal technique in the construction of models that belong to the class of rational functions, and for more general class of algebraic-function models, the Hermite-Padé-approximants are more suitable [2]. In this paper, we propose the method that is able to reduce computational cost of constructing the multivariate Padé approximants [1]. Brief description of the method is presented in the sequel.

Consider power series expansion of a multivariate function,

$$f(z) = \sum_{\nu} f_{\nu} (z - c)^{\nu}, \quad (1)$$

where $\nu = (\nu_1, \dots, \nu_d) \in \mathbb{Z}_{\geq 0}^d$ is the multi-index ($\mathbb{Z}_{\geq 0}^d$ denotes the set of non-negative integers), $z = (z_1, \dots, z_d) \in \mathbb{C}^d$ where $\mathbb{C}$ denotes the set of complex numbers, and $c = (c_1, \dots, c_d)$ is an approximation centre point. To each multi-index ν corresponds the product function $(z - c)^{\nu} = \prod_{k=1}^d (z_k - c_k)^{\nu_k}$.

The coefficients $f_{\nu} = f_{\nu_1, \dots, \nu_d}$ in (1) are computed by introducing the variable transformations $z_k = c_k + \rho_k \exp(i\theta_k)$, $k = 1, \dots, d$, $i = \sqrt{-1}$, and rewriting (1) as

$$f(\theta) = \sum_{\nu} f_{\nu} \rho^{\nu} \exp(i\nu \cdot \theta), \quad (2)$$

where $\nu \cdot \theta$ denotes the dot product so $\exp(i\nu \cdot \theta) = \prod_{k=1}^d \exp(i\nu_k \theta_k)$, $\theta = (\theta_1, \dots, \theta_d) \in \mathbb{R}^d$ ($\mathbb{R}$ is the set of real numbers), $\rho = (\rho_1, \dots, \rho_d) \in \mathbb{R}^d$ is the radius vector, and $\rho^{\nu} = \prod_{k=1}^d \rho_k^{\nu_k}$. Radius vector ρ and centre point c are problem-dependent; values are selected so that function samples are generated within the region where $f(z)$ is analytic and where the power series (1) converges.

Discretisation points of the function $f(\theta)$ on a tensor product grid, defined on the hypercube domain $\{0 \leq \theta_k \leq 2\pi, 1 \leq k \leq d\}$, can be stored as a multidimensional data array (i.e. tensor) with elements

$$F(i_1, \dots, i_d) = f(\theta_1(i_1), \dots, \theta_d(i_d)). \quad (3)$$

Abbreviated expression of (3) would be: $F(i) = f(\theta(i))$, where $i = (i_1, \dots, i_d) \in \mathbb{Z}_{>0}^d$ ($\mathbb{Z}_{>0}^d$ denotes the set of positive integers). Equidistant sampling points $\theta_k(i_k) = 2\pi(i_k - 1)/n_k$, $i_k = 1, \dots, n_k$, are used since $f(\theta)$ is a 2π -periodic function of θ . The expression (2) is the multivariate trigonometric series, therefore we use Discrete Fourier Transform (DFT) to compute the coefficients in (1):

$$f_{\nu} = \frac{1}{\rho^{\nu N}} \sum_i F(i) \exp[-i\nu \cdot \theta(i)], \quad (4)$$

the following notation is used to represent such functions: $f(\theta_1, \dots, \theta_d) = F_1(\theta_1) \cdots F_d(\theta_d)$, where $F_k(\theta_k)$ are $r_{k-1} \times r_k$ arrays with univariate functions as elements. Therefore, after discretisation, the tensor representing a semi-separable function can be written in the following matrix-product form:

$$F(i_1, \dots, i_d) = F_1(i_1) \cdots F_d(i_d), \quad (5)$$

where $F_k(i_k)$, $k = 1, \dots, d$, are index i_k - dependent matrices of the sizes $r_{k-1} \times r_k$. By exploiting the semi-separable structure, it is possible to overcome curse of dimensionality occurring when sampling on a tensor product grid. Number of stored elements in (5) is $\mathcal{O}(dnr^2)$ compared to $\mathcal{O}(n^d)$ elements in tensor product grid (3), where $n = \max(n_1, \dots, n_d)$ and $r = \max(r_0, \dots, r_d)$. Therefore, when the rank r is low, we can achieve considerable savings compared to tensor product grid. When constructing (5), we have used the tensor cross interpolation algorithm [3]. This is similar to the strategy used for the Tensor Train (TT) decomposition based on sequential application of the Singular Value Decomposition (SVD) algorithm to construct low-rank approximations of unfolding matrices [4]. However, there are two major improvements which considerably reduce number of function evaluations: a) avoid to use complete unfolding matrices (contain all tensor elements), but rather their sub-matrices which size increases during iterative process; and b) replace the SVD algorithm with the matrix cross interpolation.

When the tensor $F(i)$ is available in the matrix-product form (5), the multivariate DFT (4) can be computed efficiently as the product of univariate DFTs,

$$f_{v_1, \dots, v_d} = \prod_{k=1}^d \mathcal{F}_k(v_k) = \prod_{k=1}^d \frac{1}{n_k \rho_k^{v_k}} \sum_{i_k=1}^{N_k} F_k(i_k) \exp[-iv_k \theta_k(i_k)], \quad (6)$$

and then the power series expansion (1) is approximated as a product of $r_{k-1} \times r_k$ arrays storing univariate expansions,

$$f(z_1, \dots, z_d) = \prod_{k=1}^d \sum_{v_k} \mathcal{F}_k(v_k) (z_k - c_k)^{v_k}. \quad (7)$$

Finally, the SVD-based univariate Padé approximation algorithm [5] is applied in the summation of each univariate series in (7), and the following low-rank structure of the multivariate Padé - approximant is obtained:

$$f(z_1, \dots, z_d) = \prod_{k=1}^d \mathbb{P}_k(z_k), \quad (8)$$

where the elements of $r_{k-1} \times r_k$ arrays $\mathbb{P}_k(z_k)$, $k = 1, \dots, d$, are univariate rational functions.

Numerical examples are devised to illustrate efficiency and accuracy of the proposed method. In addition, this paper presents a practical application of the method. We show how to reconstruct solution hypersurfaces and locate singularities in the non-linear power flow problem associated with investigation of steady-state voltage stability in electric power networks [6].

References:

- [1] A. Cuyt, Multivariate Padé approximants, *Journal of Mathematical Analysis and Applications*, Vol. 96 (1983), 283-293.
- [2] M.A.H. Khan, P.G. Drazin, Y. Tourigny, The summation of series in several variable and its applications in fluid dynamics, *Fluid Dynamics Research*, Vol.33 (2003), 191-205.
- [3] D.V. Savostyanov, Quasioptimality of maximum-volume cross interpolation of tensors, *Linear Algebra and its Applications*, 432 (2014), 217-244.
- [4] I.V. Oseledets and E. E. Tyrtyshnikov, TT-cross approximation for multidimensional arrays, *Linear Algebra and its Applications*, 432 (2010), 70-88.
- [5] P. Gonnet, S. Guttel, and L.N. Trefethen, Robust Padé approximation via SVD, *SIAM Review*, Vol. 55 (2013), 101-117.
- [6] R. Živanović, Continuation via quadratic approximation to reconstruct solution branches and locate singularities in the power flow problem, *24th Mediterranean Conference on Control and Automation*, June 21-24, 2016, Greece.

Fixed and sequential designs for optimisation of fan shapes

J.-M. Azaïs¹, F. Bachoc¹, F. Gamboa¹, T. Klein¹, A. Lagnoux¹, J. Nguyen², and D. Thompson¹

¹Institut de Mathématique de Toulouse, Université Paul Sabatier, 118 route de Narbonne, 31400 Toulouse, France

²Ho Chi Minh City University of Science, VietNam

Abstract

The Institute of Mathematics of Toulouse is cooperating (among other university laboratories) with the Valeo firm on optimisation of car engine fan shapes. The study uses a complicated code to assess the performance of some geometries of fans as, for example, the max camber height, the stagger angle the chord length etc.. The optimization demands the construction of a meta-model easy to use.

In a first step a method for constructing space-filling designs of size 300- 600 in a parametric space of dimension 15-30 have been investigated. A non limitative list is given by: optimized LHS; low discrepancy sequences ; orthogonal arrays and a new method based on determinantal processes. This last method is a way of constructing repulsive point process that avoid concentration zones due to randomness. Basically the correlation function of the process is defined by a determinantal function associated to some kernel [2,3].

The comparison was conducted using classical criteria as: mindist; MST; L^2 discrepancies. We worked in very large dimension situations : 300 points in a space of dimension 15-30 is almost nothing! Results shown no clear-cut result and in particular the basic optimized LHS gave a fair trade-off between the considered criteria.

In a second step, we chose to explore sequential designs. Although we do not have any relevant model to optimize with, we know that the efficiency of the fan is a very relevant parameter and that geometries with low efficiencies are definitely not interesting. On the other hand our goal is not a simple search of a unique geometry optimizing the efficiency since the firm has to follow different specifications depending on the type of car considered.

So we decided to adopt a functional point of view considering the flow of the fan as a parameter. The maximal efficiency as a function of the flow is detected using a quadratic model and this maximal efficiency is used as a response. On this response ans using a Kriging meta-model based on Gaussian process with Matérn covariance, we used the Expected Improvement method [1,4] to add new points to a preliminary optimized LHS. In a future work this will be compared with UCB method that chose as next point the maximum of the upper limit of confidence region.

References

- [1] Clement Chevalier and David Ginsbourger. Fast Computation of the Multi-points Expected Improvement with Applications in Batch Selection. working paper or preprint, October 2012
- [2] J.B. Hough, M. Krishnapur, Y. Peres, and B. Virág. Determinantal processes and independence. *Probab. Surv.*, 3:206-229, 2006
- [3] F. Lavancier, J. Möller, and E. Rubak. Determinantal point process models and statistical inference (extended version). Technical report, Available on arxiv:1205.4818, 2014
- [4] Emmanuel Vazquez and Julien Bect. Convergence properties of the expected improvement algorithm with fixed mean and covariance functions. *J. Statist. Plann. Inference*, 140(11):3088-3095, 2010.

New Fréchet features for random distributions and associated sensitivity indices

Jean-Claude Fort^a and Thierry Klein^{b**}

April 22, 2016

^aMAP5, Université Paris Descartes, SPC, 45 rue des Saints Pères, 75006 Paris, France.

^bIMT, Université Paul Sabatier, 118 route de Narbonne, F-31062 Toulouse Cedex 9, France.

Abstract

Using contrasts we define new Fréchet features for random cumulative distribution functions. These contrasts allow to construct Wasserstein costs and our new features minimize the average costs as the Fréchet mean minimizes the mean square Wasserstein₂ distance. An example of new features is the median, and more generally the quantiles. From these definitions, we are able to define sensitivity indices when the random distribution is the output of a stochastic code. Associated to the Fréchet mean we extend the Sobol indices, and in general the indices associated to a contrast that we previously proposed.

Introduction

Nowadays the output of many computer codes is not only a real multidimensional variable but frequently a function computed on so many points that it can be considered as a functional output. This function may be the cumulative distribution function (*c.d.f.*) of a real random variable (phenomenon). Here we focused on the case of a *c.d.f.* output. To analyze such outputs one needs to choose a distance to compare various *c.d.f.*. Among the large possibilities offered by the literature we have chosen the Wasserstein distances (for more details we refer to [6]).

Thus we consider the problem of defining a generalized notion of barycenter of random probability measures on $\mathbb{R}$. It is a well known fact that the set of Radon probability measures endowed with the 2-Wasserstein distance is not an Euclidean space. Consequently, to define a notion of barycenter for random probability measures, it is natural to use the notion of Fréchet mean [4] that is an extension of the usual Euclidean barycenter. If $\mathbb{Y}$ denotes a random variable with distribution $\mathbb{P}$ taking its value in a metric space $(\mathcal{M}, d_{\mathcal{M}})$, then a Fréchet mean (not necessarily unique) of the distribution $\mathbb{P}$ is a point $m^* \in \mathcal{M}$ that is a global minimum (if any) of the functional $J(m) = \frac{1}{2} \int_{\mathcal{M}} d_{\mathcal{M}}^2(m, y) d\mathbb{P}(y)$ i.e. $m^* \in \arg \min_{m \in \mathcal{M}} J(m)$.

We present an attempt to use these tools and some extensions for analyzing computer codes outputs in a random environment, what is the subject of computer code experiments. At first we define new contrasts for random *c.d.f.* by considering generalized "Wasserstein" costs. From this, in a second step we define new features in the way of the Fréchet mean that we call Fréchet features, that we illustrate through the quantiles example. Finally we propose a sensitivity analysis of random *c.d.f.*, first from a Sobol point of view that we generalized to a contrast point of view as in [2].

1 Wasserstein distances and Wasserstein costs for unidimensional distributions

For any $p \geq 1$ we may define a Wasserstein distance between two distribution of probability, denoted F and G (their cumulative distribution functions, *c.d.f.*) on $\mathbb{R}^d$ by:

$$W_p^p(F, G) = \min_{(X, Y)} \mathbb{E} \|X - Y\|^p,$$

where the random variables (*r.v.*'s) have *c.d.f.* F and G ($X \sim F, Y \sim G$), assuming that X and Y have finite moments of order p . We call $Wasserstein_p$ space the space of all *c.d.f.* of *r.v.*'s with finite moments of order p .

As previously mentioned, in the unidimensional case where $d = 1$, it is well known that $W_p(F, G)$ is explicitly computed by:

$$W_p^p(F, G) = \int_0^1 |F^-(u) - G^-(u)|^p du = \mathbb{E} |F^-(U) - G^-(U)|^p.$$

^{**}Corresponding author. Email: jean-claude.fort@parisdescartes.fr

Here F^- and G^- are the generalized inverses of F and G that are increasing with limits 0 and 1, and U is a r.v. uniform on $[0, 1]$.

This result extends to more general contrast functions.

Definition 1.1 We call contrast functions any application c from $\mathbb{R}^2$ to $\mathbb{R}$ satisfying the "measure property" $\mathcal{P}$ defined by

$$\mathcal{P} : \forall x \leq x' \text{ and } \forall y \leq y', c(x', y') - c(x', y) - c(x, y') + c(x, y) \leq 0,$$

meaning that c defines a negative measure on $\mathbb{R}^2$.

Remark 1 If C is a convex real function then $c(x, y) = C(x - y)$ satisfies $\mathcal{P}$. This is the case of $|x - y|^p$, $p \geq 1$.

Our technical framework is the Skorohod space $\mathcal{D} := \mathcal{D}(\mathbb{R}, [0, 1])$ of all distribution functions, that is the space of all non decreasing function from $\mathbb{R}$ to $[0, 1]$ that are càd-làg with limit 0 (resp. 1) in $-\infty$ (resp. $+\infty$) equipped with the supremum norm.

Definition 1.2 (The c -Wasserstein cost) For any $F \in \mathcal{D}$, any $G \in \mathcal{D}$ and any non-negative contrast function c , we define the c -Wasserstein cost by

$$W_c(F, G) = \min_{(X \sim F, Y \sim G)} \mathbb{E}(c(X, Y)) < +\infty$$

The following theorem can be found in ([1]).

Theorem 1.1 (Cambanis, Simon, Stout [1]) Let c a function from $\mathbb{R}^2$ taking values in $\mathbb{R}$. Assume that it satisfies the "measure property" $\mathcal{P}$. Then

$$W_c(F, G) = \int_0^1 c(F^-(u), G^-(u)) du = \mathbb{E} c(F^-(U), G^-(U)),$$

where U is a random variable uniformly distributed on $[0, 1]$.

At this point we may notice that in a statistical framework many features of probability distribution can be characterized via such a contrast function. For instance an interesting case is the quantiles. Applying the remark 1 we get:

Proposition 1.1 For any $\alpha \in (0, 1)$ the contrast function (pinball function) associated to the α -quantile $c_\alpha(x, y) = (1 - \alpha)(y - x)\mathbf{1}_{x-y < 0} + \alpha(x - y)\mathbf{1}_{x-y \geq 0}$ satisfies $\mathcal{P}$.

2 Extension of the Fréchet mean to other features

A Fréchet mean $\mathcal{E}X$ of a r.v. X taking values in a metric space $(\mathcal{M}, d)$ is define as (whenever it exists):

$$\mathcal{E}X \in \operatorname{argmin}_{\theta \in \mathcal{M}} \mathbb{E} d(X, \theta)^2.$$

Thus a Fréchet mean minimizes the contrast $\mathbb{E} d(X, \theta)^2$ which is an extension of the classical contrast $\mathbb{E} \|X - \theta\|^2$ in $\mathbb{R}^d$.

Following this idea, taking c a positive contrast satisfying property $\mathcal{P}$, we define the Fréchet feature associated to c :

Definition 2.1 Assume that $\mathbb{F}$ is a random variable taking values in $\mathcal{D}$. Let c be a non negative contrast function satisfying the property $\mathcal{P}$. We define a c -contrasted feature $\mathcal{E}_c \mathbb{F}$ of $\mathbb{F}$ by:

$$\mathcal{E}_c \mathbb{F} \in \operatorname{argmin}_{G \in \mathcal{D}} \mathbb{E}(W_c(\mathbb{F}, G)).$$

This definition coincides with the Fréchet mean in the Wasserstein₂ space when using $c(F, G) = W_2^2(F, G)$.

Theorem 2.1 If c is a positive cost function satisfying the property $\mathcal{P}$, if $\mathcal{E}_c \mathbb{F}$ exists and is unique we have:

$$(\mathcal{E}_c \mathbb{F})^-(u) = \operatorname{argmin}_{s \in \mathbb{R}} \mathbb{E} c(\mathbb{F}^-(u), s).$$

For instance the Fréchet mean in the Wasserstein₂ space is the inverse of the function $u \rightarrow \mathbb{E}(\mathbb{F}^-(u))$.

Another example is the Fréchet median defines through $c = |x - y|$. Thus we consider the Wasserstein₁ space and the "contrast function" $c(F, G) = W_1(F, G)$. We obtain the Fréchet median of a random c.d.f. as :

$$(\operatorname{Med}(\mathbb{F}))^-(u) \in \operatorname{Med}(\mathbb{F}^-(u)).$$

More generally we can define an α -quantile of a random c.d.f. via the contrast function $c_\alpha(x, y)$, and we obtain $q_\alpha(\mathbb{F})$ as:

$$(q_\alpha(\mathbb{F}))^-(u) \in q_\alpha(\mathbb{F}^-(u)),$$

where $q_\alpha(X)$ is the set of the α -quantiles of the r.v. X taking its values in $\mathbb{R}$.

2.1 Sobol index

The global Sobol index quantifies the influence of the *r.v.* X_i on the output Y . This index is based on the variance (see [5]). We can define a Sobol index for the Fréchet mean of a random *c.d.f.* $\mathbb{F} = h(X_1, \dots, X_d)$. Actually we define $\text{Var}(\mathbb{F}) = \mathbb{E}W_2^2(\mathbb{F}, \mathcal{E}(\mathbb{F}))$, and

$$S_i(F) = \frac{\text{Var}(\mathbb{F}) - \mathbb{E}(\text{Var}[\mathbb{F}|X_i])}{\text{Var}\mathbb{F}}.$$

From Theorem 2.1 we get:

$$\text{Var}(\mathbb{F}) = \mathbb{E} \int_0^1 |\mathbb{F}^-(u) - \mathcal{E}(\mathbb{F})^-(u)|^2 du = \mathbb{E} \int_0^1 |\mathbb{F}^-(u) - \mathbb{E}\mathbb{F}^-(u)|^2 du = \int_0^1 \text{Var}(\mathbb{F}^-(u)) du.$$

And the Sobol index is now:

$$S_i(\mathbb{F}) = \frac{\int_0^1 \text{Var}(\mathbb{F}^-(u)) du - \int_0^1 \mathbb{E}\text{Var}[\mathbb{F}^-(u)|X_i] du}{\int_0^1 \text{Var}(\mathbb{F}^-(u)) du} = \frac{\int_0^1 \text{Var}(\mathbb{E}[\mathbb{F}^-(u)|X_i]) du}{\int_0^1 \text{Var}(\mathbb{F}^-(u)) du}.$$

2.2 Sensitivity index associated to a contrast function

The Sobol index can be extended to more general contrast functions. For a feature of a real *r.v.* associated to a contrast function c we have defined a sensitivity index (see ([2])):

$$S_{i,c} = \frac{\min_{\theta \in \mathbb{R}} \mathbb{E}c(Y, \theta) - \mathbb{E} \min_{\theta \in \mathbb{R}} \mathbb{E}[c(Y, \theta)|X_i]}{\min_{\theta \in \mathbb{R}} \mathbb{E}c(Y, \theta)}.$$

Along the same line, we now define a sensitivity index for a c -contrasted feature of a random *c.d.f.* by:

$$S_{i,c} = \frac{\min_{G \in \mathbb{W}} \mathbb{E}W_c(\mathbb{F}, G) - \mathbb{E} \min_{G \in \mathbb{W}} \mathbb{E}[W_c(\mathbb{F}, G)|X_i]}{\min_{G \in \mathbb{W}} \mathbb{E}W_c(\mathbb{F}, G)}.$$

For instance if $c = |x - y|$, $(\mathcal{E}_c \mathbb{F})^-(u)$ is the median (assumed to be unique) of the random variable $\mathbb{F}^-(u)$ and:

$$S_{i,\text{Med}} = \frac{\mathbb{E} \int_0^1 |\mathbb{F}^-(u) - \text{Med}(\mathbb{F}^-(u))| du - \mathbb{E}[\int_0^1 |\mathbb{F}^-(u) - \text{Med}[\mathbb{F}^-(u)|X_i]| du]}{\mathbb{E} \int_0^1 |\mathbb{F}^-(u) - \text{Med}(\mathbb{F}^-(u))| du}.$$

3 Conclusion

We have defined new features for a random *c.d.f.*, together with its sensitivity analysis. This theory is based on contrast functions that allow to compute Wasserstein costs. We intend to apply our methodology to an industrial problem: the PoD (Probability of Detection of a defect) in a random environment. In particular we hope that our α -quantiles will provide a relevant tool to analyze that type of data.

References

- [1] Stamatis Cambanis, Gordon Simons, and William Stout. Inequalities for $Ek(X, Y)$ when the marginals are fixed. *Z. Wahrscheinlichkeitstheorie und Verw. Gebiete*, 36(4):285–294, 1976.
- [2] J.-C. Fort, T. Klein, and N. Rachdi. New sensitivity analysis subordinated to a contrast. *Communications in Statistics-Theory and methods*, 2016.
- [3] M. Fréchet. Les éléments aléatoires de nature quelconque dans un espace distancié. *Ann. Inst. H. Poincaré, Sect. B, Prob. et Stat.*, 10:235–310, 1948.
- [4] I. M. Sobol. Sensitivity estimates for nonlinear mathematical models. *Math. Modeling Comput. Experiment*, 1(4):407–414 (1995), 1993.
- [5] C. Villani. *Topics in Optimal Transportation*. American Mathematical Society, 2003.

Comparison of Latin Hypercube and Quasi Monte Carlo Sampling Techniques

Sergei Kucherenko^a, Daniel Albrecht^b, Andrea Saltelli^c

^a*Imperial College London, London, SW7 2AZ, UK, s.kucherenko@ic.ac.uk*

^b*The European Commission, Joint Research Centre, TP 361, 21027 ISPRA(VA), ITALY*

^c*European Centre for Governance in Complexity, Universitat Autònoma de Barcelona, Catalonia, Spain*

Monte Carlo (MC) simulation employing Latin Hypercube Sampling (LHS) is one of the most popular modelling tools. While its application in areas like experimental design is well justified the efficiency of LHS in other areas such as high dimensional integration can be no better than the standard MC method based on random numbers. To provide a high efficiency of high dimensional integration high uniformity of sampling is required. LHS - being well stratified in one dimension by design, does not provide good uniformity properties in high dimensions. It is known that for high dimensional integrals the convergence rate of the MC estimates based on random sampling is $O(1/\sqrt{N})$, where N is the number of sampled points. A higher rate of convergence can be obtained by using Quasi Monte Carlo (QMC) methods based on low-discrepancy sequences. Asymptotically, QMC can provide the rate of convergence $O(1/N)$. We compare efficiencies of three sampling methods: the MC method with both random and LHS sampling, and the QMC method with sampling based on Sobol' sequences. We apply the high-dimensional Sobol' sequence generator with advanced uniformity properties (technically these are known as property **A** for all dimensions and property **A'** for adjacent dimensions). Firstly we compare L_2 discrepancies and show that the QMC method has the lowest discrepancy up to dimension 20. Secondly, we use a number of test functions of various complexities for high dimensional integration. Using global sensitivity analysis functions are classified with respect to their dependence on the input variables: functions with not equally important variables (type A), functions with equally important variables and dominant low order terms (type B) and functions with equally important variables and with dominant interaction terms (type C). Comparison shows that for types A and B functions convergence of the QMC method is close to $O(1/N)$, while the MC method has a convergence close to $O(1/\sqrt{N})$. For types C functions convergence of the QMC method significantly drops; however it still remains the most efficient method among three sampling techniques. The ANOVA decomposition in a general case can be presented as $f(x) = f_0 + \sum_i f_i(x_i) + r(x)$, where $r(x)$ are the

ANOVA terms corresponding to higher order interactions. Variance computed with the LHS design is

$$E(\varepsilon_{LHS}^2) = \frac{1}{N} \int_{H^n} [r(x)]^2 dx + O\left(\frac{1}{N}\right), \text{ and } E(\varepsilon_{MC}^2) = \frac{1}{N} \sum_i \int_{H^n} [f_i(x_i)]^2 dx + \frac{1}{N} \int_{H^n} [r(x)]^2 dx + O\left(\frac{1}{N}\right)$$

for MC. Here ε_{LHS}^2 and ε_{MC}^2 are the convergence rates.

In the ANOVA decomposition of type B functions, the effective dimension d_s is small, hence $r(x)$ is also small comparing to the main effects. In the extreme case d_s is equal to 1, and a function $f(x)$ can be presented as a sum of one-dimensional functions $f(x) = f_0 + \sum_i f_i(x_i)$. This means that only one-dimensional projections of the sampled points play a role in the function approximation. For type B functions LHS can achieve a much higher convergence rate than that of the standard MC. The results are summarized in Table 1

Table 1 Classification of functions based on the effective dimensions. Two complementary subsets of variables y and z are considered: $x = (y, z)$.

Type	Description	Relationship between S_i and S_i^{tot}	d_T	d_s	QMC is more efficient than MC	LHS is more efficient than MC
A	A few dominant variables	$S_y^{tot}/n_y \gg S_z^{tot}/n_z$	$\ll n$	$\ll n$	Yes	No
B	No unimportant subsets; only low-order interaction terms are present	$S_i \approx S_j, \forall i, j$ $S_i/S_i^{tot} \approx 1, \forall i$	$\approx n$	$\ll n$	Yes	Yes
C	No unimportant subsets; high-order interaction terms are present	$S_i \approx S_j, \forall i, j$ $S_i/S_i^{tot} \ll 1, \forall i$	$\approx n$	$\approx n$	No	No

To test the classification presented above MC, LHS and QMC integration methods were compared considering a set of test functions. Computed root mean square error (RMSE) was approximated by the formula $cN^{-\alpha}$, $0 < \alpha < 1$. For a function presented in Fig. 1 at $n=360$ the exponent for algebraic decay in the case of QMC integration $\alpha_{QMC} = 0.94$. The LHS method shows the same convergence rate $\alpha \approx 0.5$ as the MC method.

In summary QMC appears preferable to LHS overall, consideration given to the different typologies of functions. The objection that LHS can be made better by optimization (searching an optimum LHS by brute force methods) suffers from both computational hurdle and poor elegance.

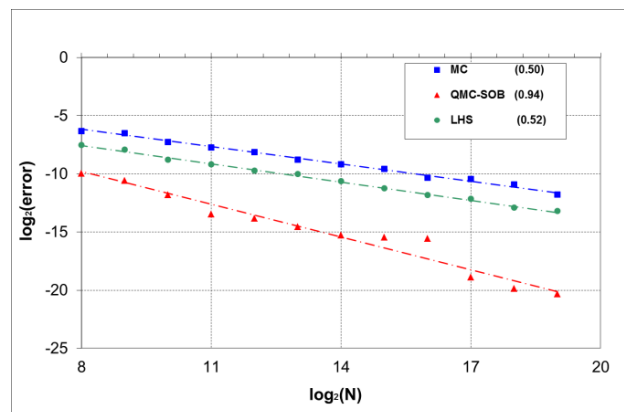


Fig. 1. RMSE versus the number of sampled points for type A model $\sum_{i=1}^n (-1)^i \prod_{j=1}^i x_j$. $n = 360$.

Global sensitivity analysis and Bayesian parameter inference for transport in a dual flowing continuum

A. Younes^{*,1,2,3}, T.A. Mara⁴, F. Delay¹.

¹ LHyGES, Université de Strasbourg/EOST, CNRS, 1 rue Blessig, 67084 Strasbourg, France.

² IRD UMR LISAH, F-92761 Montpellier, France

³ LMHE, Ecole Nationale d'Ingénieurs de Tunis, Tunisie

⁴ Université de La Réunion, PIMENT, 15 Avenue René Cassin, BP 7151, 97715 Moufia, La Réunion.

Corresponding author: younes@unistra.fr

Abstract

In this work, Global Sensitivity Analysis (GSA) and Bayesian parameter estimation are employed to interpret conservative mass transfer in a porous medium colonized by a biofilm. The GSA investigates how the model behaves with respect to its parameters as the Bayesian parameter estimation assesses the identifiability of model parameters from a set of noisy data.

GSA is useful to distinguish between influential (that contribute the most to the variability of model outputs) and non-influential parameters and hence to understand the behavior of the modeled system. In this work, GSA is performed using variance-based sensitivity indices (Sobol', 1993; Homma and Saltelli, 1996; Sobol', 2001) because they do not require any assumption regarding the linearity or the monotonicity of the model responses. These indices measure the contribution of an input (alone or via interactions with other inputs) to the output variance (Sobol', 2001). The sensitivity indices are estimated by way of a Polynomial Chaos Expansion (PCE) from a probabilistic collocation sample as in Sudret (2008) and Fajraoui et al. (2011).

Bayesian parameter estimation is performed using DREAM_(ZS) software (Laloy and Vrugt, 2012) based on the Markov Chain Monte Carlo process (MCMC). The MCMC inversion provides not only the best estimates of the parameters but also allows for exploring a large portion of parameter space in agreement with a targeted posterior distribution of the parameters. MCMC also provides the pairwise parameter correlations and the uncertainty associated with model predictions.

In the studied problem, solute transport through the porous medium colonized by biofilms is ruled by two transport equations that consider solute mass transfer in two flowing phases (porous medium and biofilm) with different velocities and dispersion coefficients (Delay et al., 2013). In a macroscopic system which cannot distinguish between the porous and the biofilm phases, the results of GSA indicate that, for weak mass transfer between phases, the output concentrations are mainly controlled by the velocity in the porous medium and by the porosity of both phases. In the case of high exchange between phases, the output concentrations are also controlled by the kinetic rate of mass transfer. The results of MCMC inversion show that transport with large mass exchange between phases is more likely subject to equifinality (i.e. lack of parameter identification) than transport with weak exchange. The Bayesian inversion also indicates that weakly sensitive parameters, such as the dispersion in each phase, can be accurately identified. However, removing these from the calibration procedures is not recommended because this model reduction might result in biased estimations of the highly sensitive parameters.

References

- Delay, F., Porel, G., Chatelier, M., 2013. Modeling by a dual-flowing continuum results of denitrification in porous media colonized by biofilms. *J. Contam. Hydrol.*, 150, 12-24.
- Fajraoui, N., F. Ramasomanana, A. Younes, T.A. Mara, P. Ackerer, and A. Guadagnini. 2011. Use of global sensitivity analysis and polynomial chaos expansion for interpretation of nonreactive transport experiments in laboratory-scale porous media. *Water Resour. Res.* 47:W02521. doi:10.1029/2010WR009639.
- Homma, T., and A. Saltelli. 1996. Importance measures in global sensitivity analysis of nonlinear models. *Reliab. Eng. Syst. Safe.* 52(1):1–17. doi:10.1016/0951-8320(96)00002-6
- Laloy, E., Vrugt, J.A., 2012. High-dimensional posterior exploration of hydrologic models using multiple-try DREAM(ZS) and high-performance computing, *Water Resour. Res.*, 48, W01526.
- Sobol', I.M. 1993. Sensitivity Estimates for Nonlinear Mathematical Models. *Mathematical Modeling and Computation.* 1(4):407-414.
- Sobol', I.M. 2001. Global sensitivity indices for nonlinear mathematical models and their Monte Carlo estimates, *Math. Comput. Simul.* 55:271-280. doi:0.1016/S0378-4754(00)00270-6
- Sudret, B. 2008. Global sensitivity analysis using polynomial chaos expansions. *Reliab. Eng. Syst. Safe.* 93(7):964-979. doi:10.1016/j.res.2006.12.019.

Real-time building design space exploration using two-sample Kolmogorov Smirnov tests to rank inputs according to multiple outputs

Torben Østergård^{a,b}, Rasmus L. Jensen^a, Steffen E. Maagaard^b

^a Aalborg University, Department of Civil Engineering, 9200 Aalborg SV. ^b MOE A/S, 8000 Aarhus, Denmark

1 Background

Building design is challenged by ever-increasing requirements towards energy demand, building costs, indoor climate, and sustainability. The industry seeks methods to support the multi-actor decision-making during the early design stage, which is characterized by an enormous design space that is difficult to model, and explore. To tackle this issue, we propose to: a) describe the variability of design parameters using uniform distributions, b) model the design space by a Monte Carlo analysis using quasi-random sampling, and c) explore the simulation results using Monte Carlo Filtering [1]. An interactive parallel coordinate plot (PCP) combined with histograms allows for real-time exploration of the simulations results by the multiple stakeholders (e.g. building owner, architects, and engineers).

Early building design typically involves many design parameters and multiple, opposing objectives (outputs). Thus, there is a need for Factor Fixing [2] of the least significant parameters prior to multi-actor meetings. Still, the PCP may contain an overwhelming number of coordinates that makes it difficult to immediately see which coordinates have been affected by a certain filter /criterion (figure 1). This study aims to see if SA using two-sample Kolmogorov-Smirnov test (KS-2-SA) can be applied to meet these challenges. KS-2-SA seems relevant since the MCF (or Factor Mapping) splits the simulations into *behavioral* and *non-behavioral* realizations [2] and MCF can be applied to models with multiple outputs (and inputs).

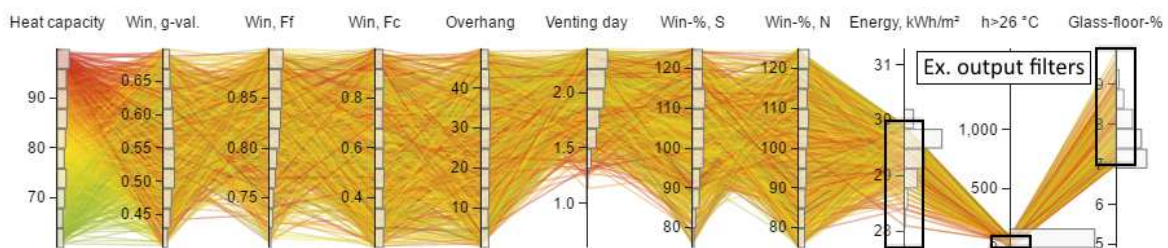


Figure 1: PCP shows input/output relationships while histograms show input/output distributions. Filters have been applied to output coordinates.

2 Methods

First, we consider sensitivity related to a single output (*Glass-floor-%*¹). Suppose that we have performed N QMC realizations of the model output. The latter is split in $J = 10$ subsamples of equal probability (i.e. each subsample is approximately of size N/J). For each subsample, MCF compares the behavioral input sample (that produced realizations of the output in the current subsample) with the non-behavioral input sample (the complementary subsample). This comparison is carried out with the two-sample Kolmogorov-Smirnov statistics D_{ij} , $j=1, \dots, J$ and $i=1, \dots, d$ with d standing for the number of input parameters (1). For each input, an average of this statistic over the number of subsamples is computed (2). We test this approach by using different sizes of subsamples J (10, 4, and 2) and different sample sizes N (200, 2.000, and 5.000).

$$(1) \quad SA_{ks2,ij} = \frac{D_{ij}}{\sum_i D_{ij}} \quad (2) \quad \overline{SA}_{ks2,i} = \frac{1}{J} \sum_j SA_{ks2,ij}$$

As a test, we first investigate how sensitivity measures from KS-2-SA compares with other sensitivity analysis methods, i.e. Pearson's R , Spearman's ρ , SRC, SRRC, Morris, and SDP (state dependent parameter SA) [3]. To compare the SA techniques, we convert the sensitivity measures into percentages (despite that Morris provides a measure for the total sensitivity whereas SRC and others estimate sensitivity from first order effects only). Secondly, we try to extend this method (of removing subsamples) to three outputs in order to enable Factor Fixing and Factor Prioritization, i.e. rank inputs in the PCP according to their influence on the three outputs simultaneously.

For this case study, we use a shoebox shaped residential building with 3.000 m² floor area. A quasi-steady state simulation model based on ISO 13790 is used to evaluate energy demand and thermal comfort (measured as the number of hours in the years during which the mean temperature exceeds 26 °C). Daylight availability is assessed by the *Glass-floor-%*. Uniform input distributions are assigned to important design variables and the

¹ The variable *Glass-floor-%* describes the amount of glazing in the building's facades measured as the percentage of glazing area per heated floor area.

combination of these constitute our “design space”. This design space is represented by up to 5.000 Monte Carlo simulations (samples) using Sobol’s (LP_T) low discrepancy sequences.

3 Findings

To evaluate the significance sample size N and subsample size J , we consider the “user-defined” output *Glass-floor-%*, which only depends on three variables: *Window frame factor*, *Window-%*, *South*, and *Window-%*, *North*. Thus, the remaining five variable inputs have no influence. To validate the results, we compare with the sensitivity percentages obtained from SRC. From figure 2.A, it seems that using subsample size of two is the best approach since the SA measures for the non-influential inputs are close to zero and to SRC measures. Figure 2.B shows how the sensitivity measure seems to improve with increasing sampling size N .

To compare different SA methods, we choose the least linear output $h > 26^\circ\text{C}$ which has a $R^2_{\text{SRC}} = 0.729$ whereas *Energy demand* and *Glass-floor-%* have values of 0.988 and 0.996. Figure 2.C shows that the sensitivity measures from KS-2 are similar to those obtained from other methods. Only the SDP approach does not match the others.

Finally, we estimate sensitivity with respect to all three variables using KS-2. We added 7 variables – all with small variations. Each output distribution is split into two subsamples which results in 8 combinations of filtering. The added variables are correctly identified as having little influence (figure 2.D). The most important inputs are: 1) *Venting day*; 2) *Win-%*, *S*; 3) *Win*, *g-value*; and 4) *Win*, *Ff*. Indeed, the histograms for these coordinates on figure 1 seem to be affected the most by the applied filters, i.e. their behavioral distributions are the “least” uniform on figure 1.

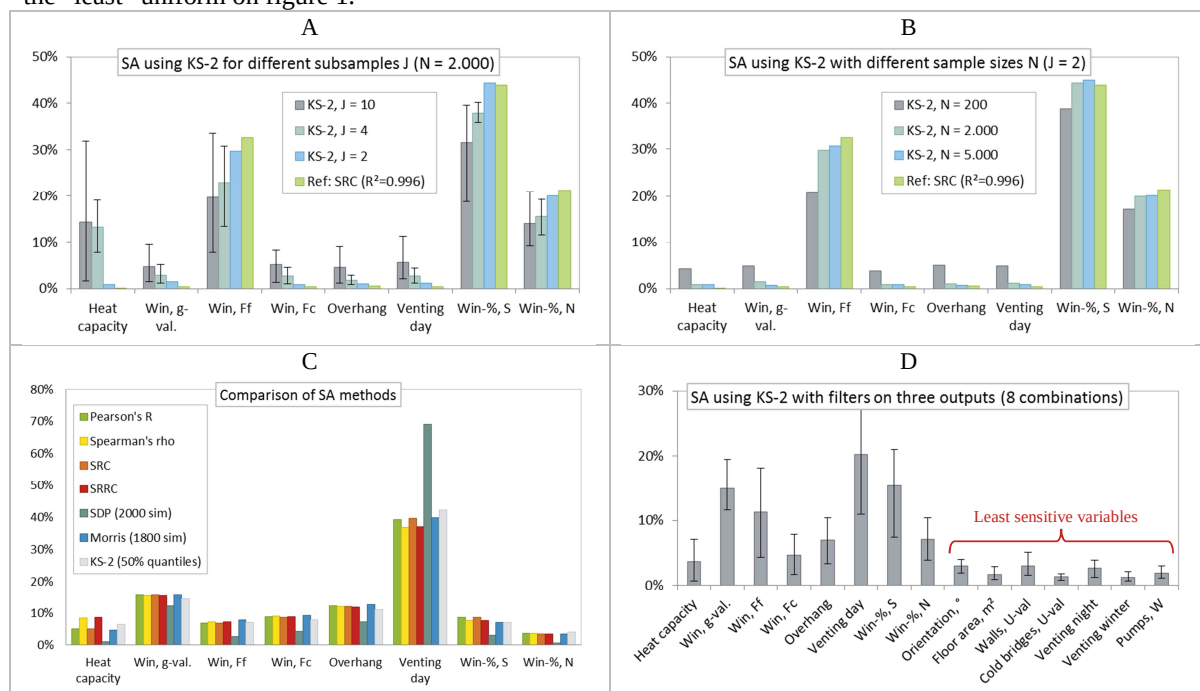


Figure 2: Comparing KS-2-test sensitivity measures when varying the quantile size (A) and sample size (B). Comparison of KS-2-test SA with other sensitivity measures based on 5.000 simulations (C). Ranking of inputs due to the combined sensitivity on 3 outputs (D).

4 Conclusion

KS-2 SA enables ranking (FP) and Factor Fixing of inputs with respect to multiple outputs. In future work, we will try to apply KS-2 SA together with PCP in real-time (Factor Mapping) so that the users immediately see which coordinates have been affected by the filtering.

5 References

- [1] MOE A/S, “Demonstration of Proactive Building Simulations,” 2016. [Online]. Available: <http://buildingdesign.moe.dk/PhD-Project/Demonstration-of-Proactive-Building-Simulations>. [Accessed: 11-Aug-2016].
- [2] A. Saltelli, M. Ratto, T. Andres, F. Campolongo, J. Cariboni, D. Gatelli, M. Saisana, and S. Tarantola, *Global sensitivity analysis: the primer*. Chichester, England: John Wiley & Sons Ltd., 2008.
- [3] M. Ratto, A. Pagano, and P. Young, “State Dependent Parameter metamodelling and sensitivity analysis,” *Comput. Phys. Commun.*, vol. 177, no. 11, pp. 863–876, 2007.

List of participants

- Ackerer Philippe
- Addi Khalid
- Azais Jean-Marc
- Becker Dirk-Alexander
- Becker William
- Benaichouche Abed
- Benoumechiara Nazih
- Bousquet Nicolas
- Bousquet Nicolas
- Browne Thomas
- Buis Samuel
- Cadéro Alice
- Carmassi Mathieu
- Damblin Guillaume
- Danko Jan
- Defaux Gilles
- Dijoux Yann
- Erlacher Christoph
- Fadji Zaouna Hassane Mamadou Maina
- Fashoto Stephen
- Fouquier Aurélie
- Iooss Bertrand
- Kucherenko Sergei
- Lamboni Matieyendou
- Lauvernet Claire
- Malet-Damour Bruno
- Mara Thierry
- Massanyi Peter
- Navarro Maria
- Oyebamiji Oluwole
- Peeters Luk
- Plischke Elmar
- Pronzato Luc
- Reiche Tatiana
- Roux Sébastien
- Ryan Edmund
- Saltelli Andrea
- Schalbart Patrick
- Song Shufang
- Spiessl Sabine M.
- Stawarz Robert
- Sudret Bruno
- Sueur Roman
- Sugimoto Takashi
- Tarantola Stefano
- Wu Qiongli
- Ostergard Torben