

A nonintrusive adaptive reduced order modeling approach for a molten salt reactor system

Alsayyari, Fahad; Tiberge, Marco; Perkó, Zoltán; Lathouwers, Danny; Kloosterman, Jan Leen

DOI

[10.1016/j.anucene.2020.107321](https://doi.org/10.1016/j.anucene.2020.107321)

Publication date

2020

Document Version

Final published version

Published in

Annals of Nuclear Energy

Citation (APA)

Alsayyari, F., Tiberge, M., Perkó, Z., Lathouwers, D., & Kloosterman, J. L. (2020). A nonintrusive adaptive reduced order modeling approach for a molten salt reactor system. *Annals of Nuclear Energy*, 141, Article 107321. <https://doi.org/10.1016/j.anucene.2020.107321>

Important note

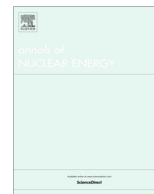
To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.



A nonintrusive adaptive reduced order modeling approach for a molten salt reactor system

Fahad Alsayyari^{*}, Marco Tiberger, Zoltán Perkó, Danny Lathouwers, Jan Leen Kloosterman

Delft University of Technology, Faculty of Applied Sciences, Department of Radiation Science and Technology, Mekelweg 15, Delft 2629JB, The Netherlands

ARTICLE INFO

Article history:

Received 8 October 2019

Received in revised form 7 January 2020

Accepted 9 January 2020

Keywords:

Reduced Order Modelling
Proper Orthogonal Decomposition
Locally adaptive sparse grids
Greedy
Data-driven
Nonintrusive
Machine learning
Uncertainty quantification
Sensitivity analysis
Molten salt reactor

ABSTRACT

We use a novel nonintrusive adaptive Reduced Order Modeling method to build a reduced model for a molten salt reactor system. Our approach is based on Proper Orthogonal Decomposition combined with locally adaptive sparse grids. Our reduced model captures the effect of 27 model parameters on k_{eff} of the system and the spatial distribution of the neutron flux and salt temperature. The reduced model was tested on 1000 random points. The maximum error in multiplication factor was found to be less than 50 pcm and the maximum L_2 error in the flux and temperature were less than 1%. Using 472 snapshots, the reduced model was able to simulate any point within the defined range faster than the high-fidelity model by a factor of 5×10^6 . We then employ the reduced model for uncertainty and sensitivity analysis of the selected parameters on k_{eff} and the maximum temperature of the system.

© 2020 The Authors. Published by Elsevier Ltd. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

1. Introduction

Complex systems such as molten salt reactors impose a modeling challenge because of the interaction between multi-physics phenomena (radiation transport, fluid dynamics and heat transfer). Such complex interaction is captured with high-fidelity, coupled models. However, these models are computationally expensive for applications of uncertainty quantification, design optimization, and control, where many repeated evaluations of the model are needed. Reduced Order Modeling (ROM) is an effective tool for such applications. This technique is based on recasting the high fidelity, high dimensional model into a simpler, low dimensional model that captures the prominent dynamics of the system with a controlled level of accuracy. Many ROM approaches can be found in literature (Antoulas et al., 2001). However, amongst studied ROM methods, Proper Orthogonal Decomposition (POD) is the suitable method for parametrized, nonlinear systems (Benner et al., 2015). The POD approach is divided into two main phases: the first is the *offline* phase, where the reduced order model is constructed by solving the high fidelity model at several points in parameter space to obtain a reduced basis space; the second is the *online*

phase, in which the reduced model is used to replace the high fidelity model in solving the system at any desired point with a reduced computational burden.

POD can be implemented intrusively by projecting the reduced basis onto the system's governing equations or non-intrusively by building a surrogate model for the POD coefficients. Many studies have successfully implemented projection based POD for nuclear applications (Buchan et al., 2013; Sartori et al., 2014; Lorenzi et al., 2017; Manthey et al., 2019; German and Ragusa, 2019).

However, for practical nuclear reactor applications, the intrusive approach is often challenging because these models are usually implemented with legacy codes that prohibit access to the governing equations, or built with coupled codes that renders modifying the governing equations a complicated task. In this case, a nonintrusive approach can be adopted to build a surrogate model for the coefficients of the POD basis. Simple interpolation or splines can be used (Ly and Tran, 2001) or for high-dimensional problems, Radial Basis Function (RBF) is usually employed (Buljak, 2011). Neural networks (Hesthaven and Ubbiali, 2018) and Gaussian regression (Nguyen and Paire, 2016) have also been studied to build the surrogate model. These approaches rely on standard sampling schemes (Monte Carlo, Latin Hypercube Sampling, tensorized uniform) to generate the snapshots. Such strategies do not take into account the dynamics of the problem and can be expensive

^{*} Corresponding author.

E-mail address: f.s.alsayyari@tudelft.nl (F. Alsayyari).

for problems parametrized on high-dimensional spaces. Audouze et al. (2009) suggested tackling this issue by combining the POD-RBF method with a greedy residual search. In this approach, the residual of the PDE is used as an error estimator by iteratively placing sampling points at locations that minimize the residual until a certain global criterion is achieved. However, this method requires repeated evaluations of the residual, which can be expensive in some solvers (e.g. matrix-free solvers) or unavailable for legacy solvers.

In this work, we propose the use of ROM method that combines the nonintrusive POD approach with the sparse grids technique (Bungartz and Griebel, 2004) to build a reduced model of a fast-spectrum molten salt system. Our approach is implemented using a previously developed algorithm (Alsayyari et al., 2019) that uses locally adaptive sparse grids as a sampling strategy for selecting the POD snapshots efficiently. The adaptivity is completely nonintrusive to the governing equations. In addition, the algorithm provides a criterion to terminate the iterations, which can be used as a heuristic estimation for the error in the developed reduced model. In this work, we extend the algorithm to deal with multiple fields of outputs. In addition, we demonstrate how local derivatives can be computed for local sensitivity analysis. The liquid-fueled system under investigation is a simplified system that captures the main characteristics of the Molten Salt Fast Reactor (Allibert et al., 2016). An in-house multi-physics tool (Tiberga et al., 2019), coupling an S_N radiation transport code with an incompressible Navier-Stokes solver, was considered as the reference model of the molten salt system. We use the developed adaptive-POD (aPOD) algorithm to construct a ROM for this reference model. We then employ the built reduced model for an uncertainty and sensitivity analysis application to study the effect of the parameters on the maximum temperature and the multiplication factor. The uncertainty and sensitivity analysis was accomplished with extensive random sampling of the reduced model. Such approach is only achievable due to the efficiency provided by the reduced model over the reference model.

The remainder of this paper is organized as follows: the POD method is briefly introduced in Section 2. Section 3 presents the sparse grids approach by introducing the interpolation technique first followed by the method for selecting the sampling points. The aPOD algorithm along with the approach to deal with multiple fields of outputs and computing the local derivatives are presented in Section 4. The model for the molten salt system is given in Section 5. The discussion of the results of constructing the reduced model is in Section 6. The uncertainty and sensitivity analysis is in Section 7. Finally, conclusions are presented in Section 8.

2. Proper orthogonal decomposition

In a nonintrusive manner, the Proper Orthogonal Decomposition can build a ROM by considering the reference, high fidelity model as a black box mapping a given input to the desired output. Let the reference model $f(y; \mathbf{x})$ be dependent on state y and a vector of input parameters $\mathbf{x}$. We can then find an expansion approximating the model as follows:

$$f(y; \mathbf{x}) \approx \sum_{i=1}^r c_i(\mathbf{x}) u_i(y), \quad (1)$$

where c_i is the expansion coefficients which depends on the input parameter $\mathbf{x}$ and $u_i(y)$ is the corresponding basis function.

The POD method seeks to find the optimal basis functions $u_i(y)$ that minimizes the error in L_2 norm,

$$\min_{u_i(y)} E = \|f(y; \mathbf{x}) - \sum_{i=1}^r c_i(\mathbf{x}) u_i(y)\|_{L_2}. \quad (2)$$

The basis functions are chosen such that they are orthonormal. Thus, the coefficients $c_i(\mathbf{x})$ can be computed as

$$c_i(\mathbf{x}) = \langle f(y; \mathbf{x}), u_i(y) \rangle, \quad (3)$$

where $\langle f(x), g(x) \rangle = \int f(x)g(x)dx$.

Assuming that the reference model is discretized ($f(y; \mathbf{x}) \rightarrow \mathbf{f}(\mathbf{x})$), the POD snapshot method finds the solution to the minimization problem using the Singular Value Decomposition (SVD). This approach begins with sampling the reference model at discrete points in parameter space $[\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_p]$, where p is the number of sampling points. Then, the corresponding outputs $[\mathbf{f}(\mathbf{x}_1), \dots, \mathbf{f}(\mathbf{x}_p)]$ can be arranged in a matrix M called the snapshot matrix. Finally, we obtain the basis vectors (also called POD modes) $\mathbf{u}_i$ as the first r left singular vectors of the SVD on the matrix M , where r is chosen to be less than or equal to the rank of the matrix M . A truncation error can be quantified using the singular values of the SVD (σ) if r is chosen to be strictly less than the rank of M ,

$$e_{tr} = \frac{\sum_{k=r+1}^n \sigma_k^2}{\sum_{k=1}^n \sigma_k^2}, \quad (4)$$

where n is the rank of M . e_{tr} quantifies the error in approximating the solutions contained in the snapshot matrix.

3. Sparse grids

For an accurate POD reduced model, the snapshots need to cover the entire dynamics of the reference model within the defined range of input parameters. Therefore, selecting an effective sampling strategy is crucial for the success of the reduced model. We propose an algorithm that is based on locally adaptive sparse grids to select the sampling points. The sparse grid algorithm builds a surrogate model for each of the POD coefficients using a Smolyak interpolant. Iteratively, the algorithm identifies a set of important points and samples their neighbouring points in the next iteration (Griebel, 1998). This process is repeated until a global convergence criterion is met. In this section we introduce the methods for the interpolation and the selection of the sampling points.

3.1. Interpolation

The Smolyak interpolation is a hierarchical interpolant that can be implemented in an iterative manner such that the accuracy is increased with each iteration (Barthelmann et al., 2000). Different basis functions can be used for the interpolant. We choose piecewise linear functions with equidistant anchor nodes since they are suitable for local adaptivity. The equidistant anchor nodes, x_j^i , corresponding to level i are defined as (Klimke, 2006)

$$m^i = \begin{cases} 1 & \text{if } i = 1, \\ 2^{i-1} + 1 & \text{if } i > 1, \end{cases} \quad (5)$$

$$x_j^i = \begin{cases} 0.5 & \text{for } j = 1 \text{ if } m^i = 1, \\ \frac{j-1}{m^i-1} & \text{for } j = 1, 2, \dots, m^i \text{ if } m^i > 1. \end{cases} \quad (6)$$

Each node defines a piecewise linear basis function ($a_{x_j^i}^i(x)$) as follows:

$$a_{x_1^1}^1 = 1 \text{ if } i = 1, \\ a_{x_j^i}^i(x) = \begin{cases} 1 - (m^i - 1)|x - x_j^i|, & \text{if } |x - x_j^i| < \frac{1}{m^i-1}, \\ 0, & \text{otherwise,} \end{cases} \quad (7)$$

The unidimensional nodes from Eq. (6) can be shown in a tree structure (Fig. 1) where the depth of the tree is assigned a level index i . The algorithm is iterative where at each iteration k , it defines a set of important points $\mathcal{Z}^k$. The criterion for selecting the important points is presented in Section 3.2. Once $\mathcal{Z}^k$ is identified, the interpolant at iteration k for a function $(c(\mathbf{x}))$ depending on a d -dimensional input x can be given by

$$A_{k,d}(c)(\mathbf{x}) = A_{k-1,d}(c)(\mathbf{x}) + \Delta A_{k,d}(c)(\mathbf{x}), \quad (8)$$

with $A_{0,d}(c)(\mathbf{x}) = 0$,

$$\Delta A_{k,d}(c)(\mathbf{x}) = \sum_{n=1}^{m_k^A} w_n^k \Theta_n(\mathbf{x}), \quad (9)$$

where m_k^A is the cardinality of $\mathcal{Z}^k$, and Θ_n is the d -variate basis function for the point $\mathbf{x}_n \in \mathcal{Z}_k$,

$$\Theta_n(\mathbf{x}) = \prod_{p=1}^d a_{x_{n,p}}^{i_p}(x_p), \quad (10)$$

where $\mathbf{x}_n$ has support nodes $(x_{n,1}^{i_1}, \dots, x_{n,d}^{i_d})$, and i_p is the level (tree depth) index for the support node $x_{n,p}^{i_p}$. w_n^k is called the *surplus* which is defined as

$$w_n^k = c(\mathbf{x}_n) - A_{k-1,d}(\mathbf{x}_n). \quad (11)$$

The union of the important points from all iterations up to k are collected in the set

$$\mathcal{X}^k = \bigcup_{l=1}^k \mathcal{Z}^l. \quad (12)$$

Because of the tree structure arrangement of the points, each point in the sparse grid $(\mathbf{x} = (x_1, \dots, x_d))$ has ancestry and descendant points. All the descendant points fall within the support of the basis function anchored at that point. The first generation descendants of a point are neighbouring points called *forward points*. The forward points for n points in the set $\mathcal{S} = \{\mathbf{x}_q | q = 1, \dots, n\}$ are defined with an operator $\Psi(\mathcal{S})$ as follows:

$$\Psi(\mathcal{S}) = \{(v_1, \dots, v_d) | \exists i, q : b(v_i) = x_{q,i} \wedge v_j = x_{q,j} \forall j \neq i, q \in [1, \dots, n], j, i \in [1, \dots, d]\}, \quad (13)$$

where $b(x)$ is a function that returns the parent of a node x from the tree. Likewise, the first generation ancestor points are called *backward points* and defined with an operator $\Psi^{-1}(\mathcal{S})$ as follows:

$$\Psi^{-1}(\mathcal{S}) = \{(v_1, \dots, v_d) | \exists i, q : b(x_{q,i}) = v_i \wedge v_j = x_{q,j} \forall j \neq i, q \in [1, \dots, n], j, i \in [1, \dots, d]\}. \quad (14)$$

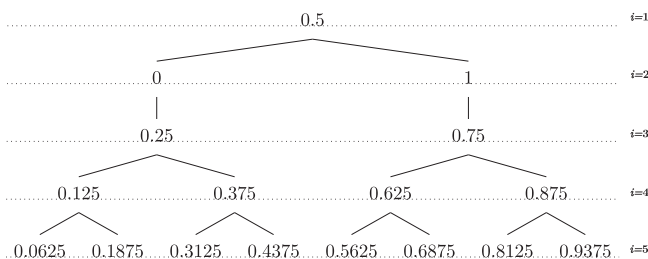


Fig. 1. Tree structure for the anchor nodes of the basis functions where the depth is assigned a level index i . At each level, nodes are added at half the distances between the nodes in the previous levels.

Finally, an operator $\Gamma(\mathcal{S})$ that return all ancestors for the points in $\mathcal{S}$ can be defined as

$$\Gamma(\mathcal{S}) = \bigcup_{l=1}^L (\Psi^{-1})^l(\mathcal{S}). \quad (15)$$

3.2. Selecting the important points

The algorithm builds the reduced model in an iterative fashion. At each iteration, we generate a set of trial points to test the model. The model is then updated according to the results of this test. Let the generated trial points be stored in the set $\mathcal{T}^k$, where k is the iteration number. The method for generating the trial points will be discussed in Section 4. For any point $\mathbf{x}_q \in \mathcal{T}^k$, we can define a local error measure ϵ_q^k in the L_2 -norm as follows:

$$\epsilon_q^k = \|\mathbf{f}(\mathbf{x}_q) - \sum_{h=1}^{r_k} A_{k,d}(c_h)(\mathbf{x}_q) \mathbf{u}_h\|_{L_2}, \quad (16)$$

where r_k is the number of POD modes selected at iteration k . The number of POD modes is selected such that the truncation error (Eq. 4) is below a defined tolerance γ_{tr} . Once ϵ_q^k is computed for all points in $\mathcal{T}^k$, we can select points with an error above a certain threshold to be stored as candidate points. The candidate points are defined as

$$\mathcal{C}^k = \{\mathbf{x}_q \in \mathcal{T}^k | \epsilon_q^k > (\gamma_{int} \|\mathbf{f}(\mathbf{x}_q)\|_{L_2} + \zeta_{abs})\}, \quad (17)$$

where γ_{int} is an interpolation threshold and ζ_{abs} is the absolute tolerance, which is introduced to deal with functions of small magnitude.

The candidate points indicate the regions in which the model needs to be enriched. To enrich the model, the ancestor points of these candidate points are first considered because ancestors have wider support. If all ancestors of the candidate points were considered important from previous iterations, that point is taken as important because the error at that point (ϵ_q^k) is above the desired threshold. This is formulated as follows:

$$\mathcal{Z}_a^k = \{\mathbf{x}_q \in \mathcal{C}^k | \Gamma(\mathbf{x}_q) \subseteq \mathcal{X}^{k-1}\}. \quad (18)$$

On the other hand, if a point $\mathbf{x}_q$ in iteration k has an error ϵ_q^k above the threshold but has also an ancestor point $\mathbf{y}_i$ which was not included in the important set in the previous iterations, $\mathbf{x}_q$ will not be marked important but its ancestor $\mathbf{y}_i$ will be marked important, because it is possible that the error ϵ_q^k was large due to missing the ancestor which has a wider support, that is

$$\mathcal{Z}_b^k = \{\mathbf{y}_i \in \Gamma(\mathbf{x}_q) | \mathbf{x}_q \in \mathcal{C}^k, \Gamma(\mathbf{x}_q) \cap \mathcal{C}^k = \emptyset \wedge \mathbf{y}_i \notin \mathcal{X}^{k-1} \wedge \Gamma(\mathbf{y}_i) \subseteq \mathcal{X}^{k-1}\}. \quad (19)$$

Then, the complete set of important points at iteration k is formed by Eqs. (18) and (19) as

$$\mathcal{Z}^k = \mathcal{Z}_a^k \cup \mathcal{Z}_b^k. \quad (20)$$

4. Algorithm

Points that are not included in the important set $\mathcal{Z}^k$ are added to the inactive set $\mathcal{I}^k$ to be tested in subsequent iterations. The trial set of the next iteration $(k+1)$ is generated as

$$\mathcal{T}^{k+1} = \left\{ \mathbf{x}_q \in \Psi(\mathcal{Z}^k) \mid \frac{\text{card}(\Psi^{-1}(\mathbf{x}_q) \cap \mathcal{X}^k)}{\text{card}(\Psi^{-1}(\mathbf{x}_q))} \geq 1 - \mu \right\} \cup \mathcal{I}^k, \quad (21)$$

where $\text{card}(\cdot)$ is the cardinality operator, and μ is a greediness parameter which has a value $\in [0, 1]$. The trial set ($\mathcal{T}^{k+1}$) is formed by the forward points of $\mathcal{Z}^k$. However, some of these forward points are excluded from being evaluated if they have some backward points not considered important in previous iterations. The number of excluded points is tuned with μ . For $\mu = 1$, all points are tested regardless of their ancestry (the algorithm in this case is more exploratory) whereas the algorithm is more efficient for $\mu = 0$ by not testing points that have any backward points not included in $\mathcal{X}^k$.

The trial set ($\mathcal{T}^{k+1}$) is then used to sample both the reduced model and the reference model to compute the error ϵ_q^{k+1} . Then, the important points ($\mathcal{Z}^{k+1}$) are identified and added to the snapshot matrix. Each update to the snapshot matrix generates a complete new set of POD modes, which requires recomputing the interpolant $A_{k,d}(c)(\mathbf{x})$ because of its dependence on the POD modes. Specifically, the surpluses ($w_{q,h}^k$) corresponding to POD mode $\mathbf{u}_h$ need to be recomputed with each POD update. The surpluses are just the deviations of the interpolant from the true value. Therefore, an easy way to update the surpluses after each iteration is as follows:

$$\hat{w}_{q,g}^k = \sum_{h=1}^{r_k} w_{q,h}^k < \mathbf{u}_h, \hat{\mathbf{u}}_g > \quad g = 1, \dots, r_{k+1}, \quad (22)$$

where $\hat{\mathbf{u}}_g$ is the g th POD mode after updating the snapshot matrix, $\mathbf{u}_h$ is the h th POD mode before updating the snapshot matrix, $w_{q,h}^k$ is the surplus at iteration k corresponding to the point $\mathbf{x}_q \in \mathcal{X}^k$ and POD mode $\mathbf{u}_h$, and $\hat{w}_{q,g}^k$ is the updated surplus corresponding to $\mathbf{x}_q \in \mathcal{X}^k$ and $\hat{\mathbf{u}}_g$. For further reading regarding the adaptive sparse grids technique and the derivation of Equation 22, see Alsayyari et al. (2019) and the references within. Fig. 2 summarizes the algorithm.

4.1. Multiple outputs

To deal with models of multiple outputs, we can build a different ROM model for each output, which entails running the adaptive-POD algorithm separately for each output. With such an approach, managing the output field data is important to prevent multiple costly evaluations of the same point. This can be achieved by storing all output fields for any full model evaluation in a data bank, which the algorithm is directed to access when a point is required more than once in different output field constructions. With this strategy, the separate runs of the algorithm are

performed in series rather than parallel in order to avoid full evaluations of the same point. Another approach is to combine the output fields by stacking them into a composite vector which is then treated as a single output in the snapshot matrix. In this approach, only a single ROM is built to represent all outputs. Since the first approach is a straightforward application of the algorithm, in this section, we show how the second approach is implemented.

Let the outputs be represented by $\mathbf{f}_1(\mathbf{x}), \dots, \mathbf{f}_o(\mathbf{x})$ where o is the number of output fields. The snapshot matrix is formed by stacking the output fields as

$$[(\mathbf{f}_1^T(\mathbf{x}_1), \dots, \mathbf{f}_o^T(\mathbf{x}_1))^T, \dots, (\mathbf{f}_1^T(\mathbf{x}_p), \dots, \mathbf{f}_o^T(\mathbf{x}_p))^T]. \quad (23)$$

We can compute the local error measure (Eq. (16)) in each output $\mathbf{f}_s(\mathbf{x}_q)$ separately

$$\epsilon_{s,q}^k = \|\mathbf{f}_s(\mathbf{x}_q) - \sum_{h=1}^{r_k} A_{k,d}(c_h)(\mathbf{x}_q) \mathbf{u}_{s,h}\|_{L_2}. \quad (24)$$

Different interpolation thresholds and absolute tolerances can be defined for each output. A point $\mathbf{x}_q$ is admitted to the candidate set (Eq. (17)) if the corresponding error $\epsilon_{s,q}^k$ at any of the output fields ($s \in [1, \dots, o]$) is greater than the defined threshold

$$\begin{aligned} \mathcal{C}^k &= \{\mathbf{x}_q \in \mathcal{T}^k \mid \exists s = [1, \dots, o] : \epsilon_{s,q}^k \\ &> \gamma_{\text{int},s} \|\mathbf{f}_s(\mathbf{x}_q)\|_{L_2} + \zeta_{\text{abs},s}\}, \end{aligned} \quad (25)$$

where $\gamma_{\text{int},s}$ and $\zeta_{\text{abs},s}$ are respectively the interpolation threshold and the absolute tolerance defined for output $\mathbf{f}_s(\mathbf{x})$.

The algorithm is terminated when a global criterion is met. We define this criterion to be

$$\epsilon_{s,q}^k < (\zeta_{\text{rel},s} \|\mathbf{f}_s(\mathbf{x}_q)\|_{L_2} + \zeta_{\text{abs},s}), \quad \forall \mathbf{x}_q \in \mathcal{T}^k, \quad s = 1, \dots, o, \quad (26)$$

where $\zeta_{\text{rel},s}$ is the global relative tolerance set for output $\mathbf{f}_s(\mathbf{x})$. Note that the multiple-outputs approach can yield a different performance compared to the single-output approach in terms of points selected for evaluations. This is because the POD basis is constructed differently. In the single-output approach, the POD modes are tailored to that output specifically whereas in the multiple-outputs approach the POD modes contain information for all output fields.

4.2. Calculation of local sensitivities

To compute local sensitivities, we can find an analytical expression for the derivatives of each output with respect to the inputs. The derivative of the ROM model in Eq. (1) with respect to the g th dimension x_g is

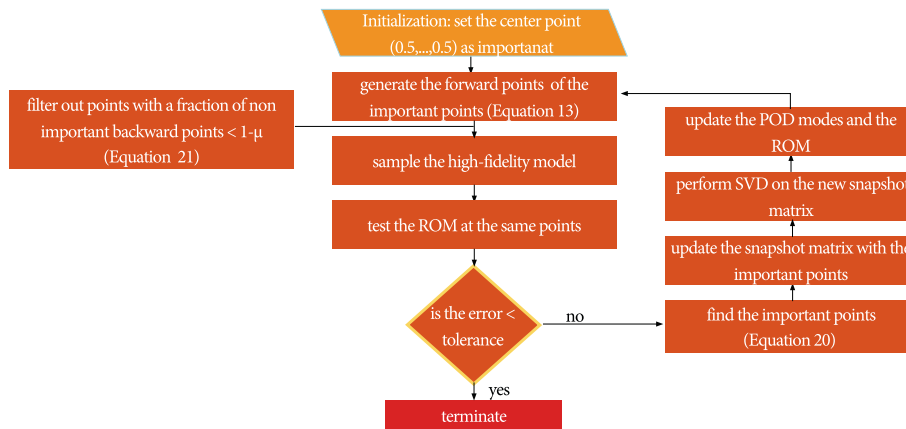


Fig. 2. A graphical scheme of the adaptive-POD (aPOD) algorithm..

$$\frac{\partial}{\partial \mathbf{x}_g} \mathbf{f}(\mathbf{x}) = \frac{\partial}{\partial \mathbf{x}_g} \sum_{i=1}^r c_i(\mathbf{x}) \mathbf{u}_i. \quad (27)$$

The ROM model interpolates $c_i(\mathbf{x})$ with the operator $A_{k,d}(c)(\mathbf{x})$. Using Eqs. (8) and (9), Eq. (27) becomes

$$\frac{\partial}{\partial \mathbf{x}_g} \mathbf{f}(\mathbf{x}) = \frac{\partial}{\partial \mathbf{x}_g} \sum_{i=1}^r \left(\sum_{n=1}^{m_k^A} w_{n,i}^k \Theta_n(\mathbf{x}) \right) \mathbf{u}_i, \quad (28)$$

$$= \sum_{i=1}^r \mathbf{u}_i \sum_{n=1}^{m_k^A} w_{n,i}^k \frac{\partial}{\partial \mathbf{x}_g} \prod_{p=1}^d a_{x_{n,p}}^{i_p}(x_p), \quad (29)$$

$$= \sum_{i=1}^r \mathbf{u}_i \sum_{n=1}^{m_k^A} w_{n,i}^k \left[\frac{\partial}{\partial \mathbf{x}_g} a_{x_{n,g}}^{i_g}(x_g) \right] \prod_{p \neq g}^d a_{x_{n,p}}^{i_p}(x_p), \quad (30)$$

where the derivative of the unidimensional basis function $\frac{\partial}{\partial x} a_{x_n}^i(x)$ (dropping the dependence on the dimension g) is computed as

$$\frac{\partial}{\partial x} a_{x_n}^i = 0 \quad \text{if } i = 1, \\ \frac{\partial}{\partial x} a_{x_n}^i(x) = \begin{cases} -(m^i - 1) \frac{x - x_n^i}{|x - x_n^i|}, & \text{if } |x - x_n^i| < \frac{1}{m^i - 1}, \quad x \neq x_n^i \\ 0, & \text{if } |x - x_n^i| \geq \frac{1}{m^i - 1} \\ \text{Not defined,} & \text{if } x = x_n^i, \end{cases} \quad (31)$$

It is evident that due to the choice of piecewise linear basis functions, our reduced model is non-differentiable at the anchor nodes x_n^i , which implies that we cannot compute local derivatives at the sampled snapshots, including the nominal point. However, we can compute the local derivatives at two points very close to the nominal values and average them out to have a measure of the local sensitivities at the nominal point.

5. Molten salt system

In this work, we construct a reduced order model of a simplified system representative of the main characteristics of the Molten Salt Fast Reactor (Allibert et al., 2016): strong coupling between neutronics and thermal-hydraulics, fast spectrum, and transport of precursors. The problem was developed as a benchmark for multi-physics tools dedicated to liquid-fuel fast reactors (Aufiero and Rubiolo, 2018; A. Laureau et al., 2015).

Fig. 3 depicts the problem domain: a 2 m side square, 2-dimensional cavity filled with fluoride molten salt at initial temperature of 900 K. The cavity is surrounded by vacuum and insulated; salt cooling is simulated via a heat sink equal to $h(T_{\text{ext}} - T)$, where $T_{\text{ext}} = 900$ K and h is a volumetric heat transfer coefficient. Zero-velocity boundary conditions are applied to all walls except the top lid, which moves at $v_{\text{lid}} = 0.5 \text{ ms}^{-1}$. The steady-state solution is sought with criticality eigenvalue calculations normalizing the reactor power to P_0 . Fluid properties are con-

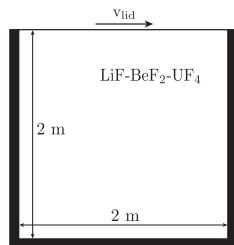


Fig. 3. Simplified molten salt fast system: square cavity domain. It is insulated, surrounded by vacuum, and filled with molten fluoride salt at initial temperature of 900 K. The top lid moves with velocity $v_{\text{lid}} = 0.5 \text{ ms}^{-1}$.

stant with temperature and uniform in space. Neutronics data are condensed into 6 energy groups and temperature corrected only via density feedback, to avoid the complexities related to Doppler feedback modeling; delayed neutron precursors are divided into 8 families. The flow is laminar and buoyancy effects are modeled via the Boussinesq approximation. Cross sections are corrected according to

$$\Sigma(T) = \Sigma(T_{\text{ref}}) \frac{\rho(T)}{\rho(T_{\text{ref}})} = \Sigma(T_{\text{ref}})(1 - \beta_{\text{th}}(T - T_{\text{ref}})) \quad (32)$$

where $T_{\text{ref}} = 900$ K and $\rho(T_{\text{ref}})$ is the density at which macroscopic cross sections are provided. They correspond to the reference values chosen for the Boussinesq approximation. β_{th} is the thermal expansion coefficient. We refer to Aufiero and Rubiolo (2018) and A. Laureau et al. (2015) for a more detailed description of the problem.

An in-house multi-physics tool is used to model the molten salt system. It couples a solver for the incompressible Navier-Stokes equations (DGFlows) with a neutronics code solving the multi-group S_N Boltzmann equation coupled with the transport equations for the delayed neutron precursors (PHANTOM - S_N). Both codes are based on the Discontinuous Galerkin Finite Element method for space discretization. Fig. 4 displays the structure of the multi-physics tool and the data exchanged between the codes. The average temperature on each element (T_{avg}) is outputted to PHANTOM - S_N , which applies the density feedback on cross sections taken from the library at 900 K, according to Equation 32. Then, the neutronics problem is solved taking the velocity field ($\mathbf{u}$) from DGFlows as another input for the delayed neutron precursors equation. Finally, the fission power density (P_{fiss}) is transferred to the CFD code. The steady state solution is sought by iterating DGFlows and PHANTOM - S_N until convergence. More details on the multi-physics tool can be found in Tiberge et al. (2019).

Simulations of the molten salt system were performed choosing a 50×50 uniform structured mesh, with a second-order polynomial discretization for the velocity and a first-order one for all the other quantities. An S_2 discretization was chosen for the angular variable. Fig. 5 shows the steady state fields (velocity magnitude, temperature, and total flux) obtained for the nominal values of the input parameters. The nominal multiplication factor in this configuration is $k_{\text{eff}} = 0.99295$. The upper bounds for each of the six energy groups are shown in Table 1 along with the space averaged flux (Φ_{avg}) for each group in the nominal case.

6. Results

A ROM model was built for the molten salt system by considering 27 input parameters. We assumed a uniform distribution for all parameters. The parameters and the corresponding percentage variation from the nominal values are summarized in Table 2, where P_0 is the initial power, β_{th} is the thermal expansion coefficient, $\Sigma_{f,g}$ is the fission cross section for group g , β_i is the delayed neutron fraction for precursors family i , λ_i is the decay constant for precursors

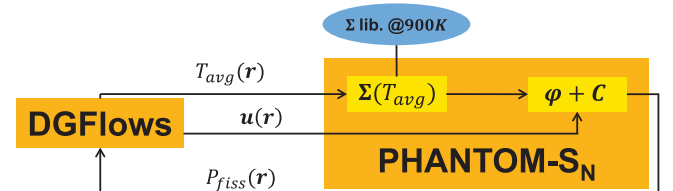


Fig. 4. Computational scheme of the multi-physics tool representing the high-fidelity model. The CFD code, DGFlows, exchanges data with the radiation transport code, PHANTOM - S_N , at each iteration due to the coupling between the physics characterizing the molten salt nuclear system.

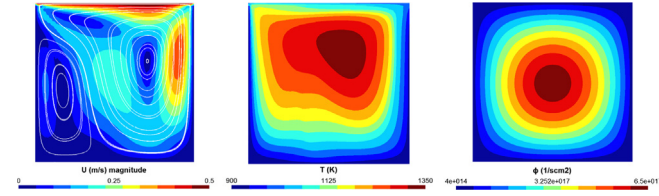


Fig. 5. Velocity magnitude, temperature, and total flux fields representing the steady state solution of the simplified MSFR problem for nominal values of the input parameters.

family i , v_{lid} is the lid velocity, ν is the viscosity, and h is the heat transfer coefficient. Since we aim at using the reduced model for uncertainty and sensitivity analysis, we assigned a variation of $\pm 10\%$ for parameters with typical experimental uncertainties whereas we vary design parameters (P_0 , v_{lid} and h) by $\pm 20\%$. Our interest is in the effect of these parameters on the spatial distribution of the total flux $\Phi(\mathbf{r})$, the temperature $T(\mathbf{r})$, and the value of the effective multiplication factor k_{eff} . Therefore, the reference model has 27 inputs and returns a value for the k_{eff} and two field vectors each of length 7500 corresponding to the coefficients of the discontinuous Galerkin expansion for the total flux Φ and temperature T . In this work, we compare the stacking of the outputs approach described in Subsection 4.1 with the single-output approach. For the multiple-outputs approach, the snapshot matrix for the outputs evaluated at points $[\mathbf{x}_1, \dots, \mathbf{x}_p]$ is computed as $\left[\left(\Phi_1^T, T_1^T, k_{\text{eff},1} \right)^T, \dots, \left(\Phi_p^T, T_p^T, k_{\text{eff},p} \right)^T \right]$.

The global relative tolerances ζ_{rel} for Φ and T were set to be 10^{-2} , which means we require the error in the L_2 norm for these fields to be less than 1%. For k_{eff} , we require the error to be less than 50 pcm, so we set ζ_{rel} for k_{eff} to be 50×10^{-5} . The interpolation threshold (γ_{int}) was chosen to be one order of magnitude less than the set relative tolerances. Therefore, γ_{int} was 10^{-3} for both Φ and T and was set to be 5×10^{-5} for k_{eff} .

We first built a reduced model using a greediness value $\mu = 1$. For the multiple-outputs approach, the algorithm required 4495 reference model evaluation to converge. However, only 142 points were included in the important set. The small number of selected important points is an indication of oversampling. The algorithm was then run again with $\mu = 0$. In this case, the algorithm sampled 472 points with 105 important points included in the snapshot matrix, which is a reduction by about a factor of 10 in the number of evaluations compared with the case of $\mu = 1$. Each reference model evaluation takes about 1.5 h to run (performed on a Linux cluster using 1 CPU operating at 2.60 GHz). Therefore, this reduction in number of evaluations is massive in computational time. In order to test the model, 1000 Latin Hypercube Sampling (LHS) points were generated. LHS is a method to generate unbiased random points in higher dimensional spaces by partitioning the hypercube first. Then, drawing one sample from each partition. These generated points were not part of the snapshot matrix. Note that the reduced model was trained only on the important set. The rest of the model evaluations served as trial points but were not included in the snapshot matrix. In machine learning terminology, the important set is the training set and the rest of the evaluations

served the function of the validation set (Ripley, 1996). Therefore, the generated 1000 unbiased random points in the test set represent 10 times more testing points than training points. Running the reduced model on the 1000 testing points needed only about one second on a personal computer.

Table 3 summarizes the maximum L_2 norm error found for each output. It is evident that all tested points resulted in errors well below the set tolerances. We also compare the results of the single-output approach to the multiple-outputs approach in the same table. While both approaches satisfied the required tolerances, the number of full model evaluations required in the offline stage was different. The single-output approach required fewer evaluations compared to the multiple-outputs approach. This is due to the fact that the POD modes in the single-output approach are tailored to that output field. The algorithm in this case, samples points to construct a specific reduced model satisfying the desired tolerance for that output. In the multiple-outputs approach, on the other hand, the algorithm uses POD modes containing information for all output fields, which require more points to satisfy the desired tolerances for every output fields. However, because the reduced model is enriched with every additional sampling point, the multiple-outputs model has a slightly less error in the online phase compared to the single-output approach.

Fig. 6 shows the distribution of the L_2 norm error for the tested 1000 random points for each output in the reduced model of the multiple-outputs approach and $\mu = 0$. A comparison between the temperature distributions of the reduced model and the reference full order model at the point that resulted in the maximum error is shown in Fig. 7. The L_2 norm error for this case was 0.2% while the maximum absolute difference locally was 13.9 K, which is about 1% of the maximum local temperature (about 1482.6 K). Both cases of $\mu = 1$ and $\mu = 0$ converged with 3 iterations ($k = 3$). To highlight the cost effectiveness of the adaptive approach, for such 27-dimensional problem, the classical (non-adaptive) sparse grid approach would require 27829 points after 3 iterations, which is extremely expensive to run.

Table 4 summarizes the number of unique nodes per dimension, which was found to be the same for both the single and multiple-outputs approaches. This number is indicative of the linearity/non-linearity of the reference model. During the construction stage, the algorithm captures the degree of linearity of the output of the reference model with respect to each dimension within the defined range. A value of 3 means that the algorithm considered that dimension to be constant because after building a constant interpolant at the root 0.5, the error in the model was found to be within the defined tolerances at the children points $\{0, 1\}$. The algorithm then stopped further refinements along that dimension. A value of 5 indicates that the model is piecewise linear in the segments $(0, 0.5)$ and $(0.5, 1)$ with respect to that dimension because the refinement is stopped after testing the piecewise linear interpolant using the first 3 points $\{0.5, 0, 1\}$ at the children $\{0.25, 0.75\}$. A value higher than 5 indicates that the model is non-linear along that dimension.

It is evident from the number of unique nodes that the algorithm found the outputs of the model to be constant (within the set tolerances) with respect to β_i and λ_i , which means varying these parameters within the 10% range does not significantly affect the defined outputs. Additionally, the model was found to be piecewise

Table 1
Average group flux in the nominal case along with the upper energy bound for each group.

Energy group	1	2	3	4	5	6
Upper bound [keV]	20000	2231	497.9	2.479	5.531	0.7485
$\Phi_{\text{avg}} [\text{cm}^{-2} \text{s}^{-1}] \times 10^{16}$	1.22	4.92	9.19	5.94	4.74	1.43

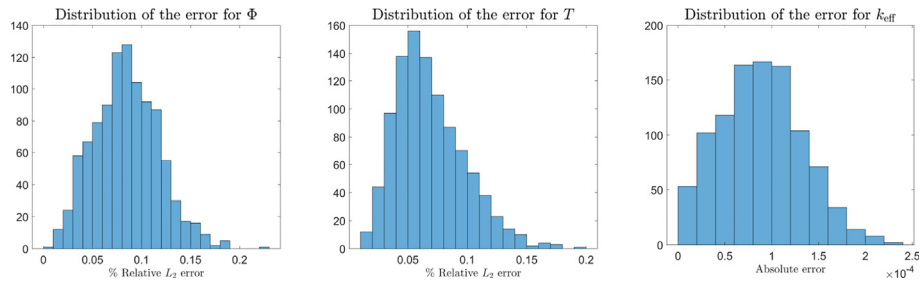
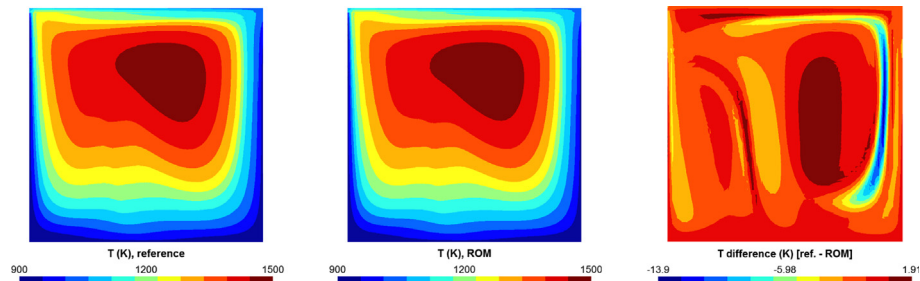
Table 2

Nominal values and the corresponding variation for the considered parameters.

Parameter	Nominal value	Percentage variation	Parameter	Nominal value	Percentage variation
P_0 [W]	10^9	$\pm 20\%$	β_7	6.05×10^4	$\pm 10\%$
β_{th} [K $^{-1}$]	2×10^{-4}	$\pm 10\%$	β_8	1.66×10^{-4}	$\pm 10\%$
$\Sigma_{f,1}$ [cm $^{-1}$]	1.11×10^{-3}	$\pm 10\%$	λ_1 [s $^{-1}$]	1.25×10^{-2}	$\pm 10\%$
$\Sigma_{f,2}$ [cm $^{-1}$]	1.08×10^{-3}	$\pm 10\%$	λ_2 [s $^{-1}$]	2.83×10^{-2}	$\pm 10\%$
$\Sigma_{f,3}$ [cm $^{-1}$]	1.52×10^{-3}	$\pm 10\%$	λ_3 [s $^{-1}$]	4.25×10^{-2}	$\pm 10\%$
$\Sigma_{f,4}$ [cm $^{-1}$]	2.58×10^{-3}	$\pm 10\%$	λ_4 [s $^{-1}$]	1.33×10^{-1}	$\pm 10\%$
$\Sigma_{f,5}$ [cm $^{-1}$]	5.36×10^{-3}	$\pm 10\%$	λ_5 [s $^{-1}$]	2.92×10^{-1}	$\pm 10\%$
$\Sigma_{f,6}$ [cm $^{-1}$]	1.44×10^{-2}	$\pm 10\%$	λ_6 [s $^{-1}$]	6.66×10^{-1}	$\pm 10\%$
β_1	2.33×10^{-4}	$\pm 10\%$	λ_7 [s $^{-1}$]	1.63	$\pm 10\%$
β_2	1.03×10^{-3}	$\pm 10\%$	λ_8 [s $^{-1}$]	3.55	$\pm 10\%$
β_3	6.81×10^{-4}	$\pm 10\%$	v_{lid} [m/s]	0.5	$\pm 20\%$
β_4	1.37×10^{-3}	$\pm 10\%$	ν [m 2 /s]	0.025	$\pm 10\%$
β_5	2.14×10^{-3}	$\pm 10\%$	h [W/m 2 K]	10^6	$\pm 20\%$
β_6	6.41×10^{-4}	$\pm 10\%$			

Table 3Maximum L_2 error in each output with respect to the reference model after testing the reduced model on 1000 random points. The total number of full model evaluations in the offline stage for each ROM construction is also shown.

		Φ	T	k_{eff}	Total number of evaluations
Multiple outputs	$\mu = 1$	0.18%	0.14%	23 pcm	4495
	$\mu = 0$	0.22%	0.20%	22 pcm	472
Single output	$\mu = 1$	0.35%	0.14%	23 pcm	3548
	$\mu = 0$	0.35%	0.25%	33 pcm	348

**Fig. 6.** Histogram showing the error in each of the outputs resulting from testing the reduced model on 1000 random points.**Fig. 7.** Temperature distribution at the point of maximum error showing the reference model (left), the ROM model (center), and the distribution of the difference (right). Note the change of the colour bar scale in the difference plot (right).

linear with respect to the power, velocity, thermal expansion coefficient, viscosity, and the fission cross section for the groups 1–4. However, for the lowest energy groups (group 5 and 6), the model was nonlinear. This can be explained by the fact that the flux distributions for all groups were not changing significantly due to the homogeneity of the changes to the system. In addition, the group fluxes were found to have the same order of magnitude as shown in Table 1 for the nominal case. However, the nominal values of the

fission cross section for $\Sigma_{f,5}$ and $\Sigma_{f,6}$ are higher compared to the other fast groups, which weigh more in the calculation of k_{eff} . By examining the cause for the additional unique points along $\Sigma_{f,5}$ and $\Sigma_{f,6}$, we found that they were triggered purely by k_{eff} and not by Φ or T . The model was also nonlinear in the heat transfer coefficient. The negligible effect of β_i and λ_i explains the reason for the massive reduction in number of evaluations with the setting $\mu = 0$. The algorithm in this case recognized that β_i and λ_i have

Table 4
Number of unique nodes per dimension.

Parameter	number of unique nodes	Parameter	number of unique nodes
P_0	5	β_7	3
β_{th}	5	β_8	3
$\Sigma_{f,1}$	5	λ_1	3
$\Sigma_{f,2}$	5	λ_2	3
$\Sigma_{f,3}$	5	λ_3	3
$\Sigma_{f,4}$	5	λ_4	3
$\Sigma_{f,5}$	9	λ_5	3
$\Sigma_{f,6}$	9	λ_6	3
β_1	3	λ_7	3
β_2	3	λ_8	3
β_3	3	v_{lid}	5
β_4	3	v	5
β_5	3	h	9
β_6	3		

no effect within the defined range and stopped sampling points along these dimensions. Since β_i and λ_i amount to 16 out of the 27 dimensions, the reduction in number of points was massive.

7. Uncertainty quantification and sensitivity analysis

In this section, we demonstrate the potential of the built ROM model in an application of uncertainty quantification and sensitivity analysis. We study the effect of the selected input parameters on the maximum temperature and the multiplication factor k_{eff} . The resulting ROM can be sampled cheaply at any point within the specified range. The ROM model from the multiple-outputs approach and $\mu = 0$ is employed for the study in this section. However, we do not expect differences in the results if any of the other 3 ROM models developed in Section 6 were used instead. We use Latin Hypercube Sampling to sample the reduced model with 100,000 random points. The density histograms approximating the Probability Distribution Function (PDF) are shown in Fig. 8. For comparison, the densities resulting from running the reference model on the 1000 testing points are also shown in the figure. The density histogram shows a distribution close to a normal distribution, which can be explained by the fact that all input parameters are assumed to have uniform distribution and the model is linear or almost linear in these parameters. Therefore, the sum of these uniform distribution approaches the normal distribution. The normal probability plot in Fig. 9 confirms that the distribution is normal within the middle range while the deviation from the normal

is seen at the tails of the distribution. The mean of the maximum temperature was found to be at 1336.5 K with standard deviation equal to 61.1 K while the mean of k_{eff} was 0.99229 with standard deviation equal to 0.016.

Local and global sensitivity analyses were also performed using the built ROM. For the local sensitivities, Table 5 presents the averaged derivatives computed from several points within a distance of 10^{-14} (measured in the unit hypercube $[0, 1]^d$) from the input's nominal values. In order to provide a better comparison of the effect of the parameters, the computed derivatives in the table are normalized by the ratio $R_0/x_{p,0}$, where R_0 is the desired response (maximum temperature or k_{eff}) computed at the nominal values of the input parameters $x_{p,0}$.

The results show that the maximum temperature is mainly affected by the initial power P_0 and the heat transfer coefficient h . This is expected because these two parameters directly control the amount of energy present in the system. Higher initial power increases the amount of energy in the system which directly raises the temperature. The heat transfer coefficient, on the other hand, is negatively correlated with T_{max} because lower h decreases the amount of energy being extracted from the system causing the temperature to rise.

The thermal expansion coefficient is related to the natural convection phenomenon. Forced and natural convection play a competing role in terms of mixing of the salt in the cavity. There are two vortices in the cavity as shown by the streamlines in Figure 5 (left) for the nominal case. When forced convection increases, the larger vortex grows causing the vortex centre to move towards the cavity centre. In this case, salt in the central region of the cavity would always circulate around the centre where the fission power is maximum. On the other hand, when natural convection increases, the smaller vortex in the bottom left corner becomes larger causing the salt to pass through the centre then transported close to the boundaries of the cavity where the thermal energy is minimum. Hence, in the range of variations considered in this work, natural convection tends to redistribute the heat in the cavity, whereas forced convection has the opposite effect. Higher β_{th} causes natural convection to be more prevalent over forced convection. This causes the temperature to be more uniform. For this reason, β_{th} is negatively correlated with T_{max} . The viscosity, on the other hand, has the opposite effect. Increasing the viscosity reduces the mixing of the liquid, which creates more concentrated hot spots that increase the maximum temperature. The lid velocity is also positively correlated with the maximum temperature

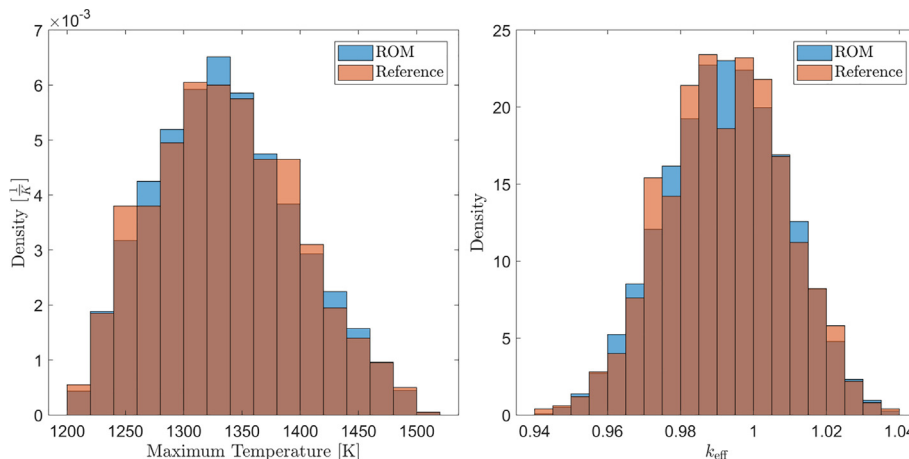


Fig. 8. Density histograms of the maximum temperature (left) and the multiplication factor k_{eff} (right) by sampling the reduced model with 100,000 points. The distributions of same variables from sampling the reference model with the 1000 testing points are also shown. Note that the histogram is normalized such that the sum of the areas of the bars equals to 1..

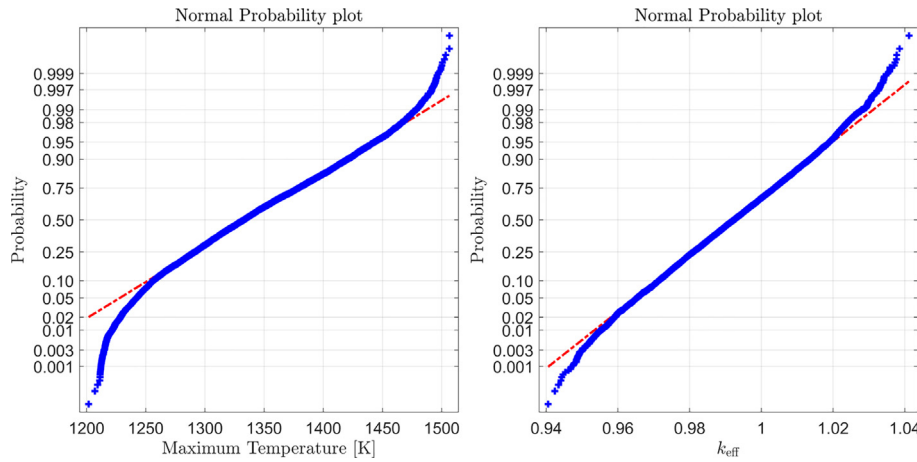


Fig. 9. Normal probability plots for the maximum temperature (left) and the multiplication factor k_{eff} (right) showing the distribution to be normal within the middle parts but deviating from the normal distribution at the tails.

Table 5

Normalized local sensitivities of the maximum temperature (T_{max}) and k_{eff} with respect to the parameters around the nominal values. The derivatives are normalized by the ratio of output nominal value ($T_{\text{max},0}$ and $k_{\text{eff},0}$) to the nominal values of the input parameters $x_{p,0}$.

x_p	$\frac{\partial T_{\text{max}}}{\partial x_p} \frac{T_{\text{max},0}}{x_{p,0}}$	$\frac{\partial k_{\text{eff}}}{\partial x_p} \frac{k_{\text{eff},0}}{x_{p,0}}$	x_p	$\frac{\partial T_{\text{max}}}{\partial x_p} \frac{T_{\text{max},0}}{x_{p,0}}$	$\frac{\partial k_{\text{eff}}}{\partial x_p} \frac{k_{\text{eff},0}}{x_{p,0}}$
P_0	0.289	-0.012	β_7	0	0
β_{th}	-0.036	-0.012	β_8	0	0
$\Sigma_{f,1}$	-2×10^{-5}	0.012	λ_1	0	0
$\Sigma_{f,2}$	-8×10^{-5}	0.041	λ_2	0	0
$\Sigma_{f,3}$	-6×10^{-5}	0.101	λ_3	0	0
$\Sigma_{f,4}$	-2×10^{-5}	0.11	λ_4	0	0
$\Sigma_{f,5}$	3×10^{-5}	0.182	λ_5	0	0
$\Sigma_{f,6}$	9×10^{-5}	0.145	λ_6	0	0
β_1	0	0	λ_7	0	0
β_2	0	0	λ_8	0	0
β_3	0	0	v_{lid}	0.0003	10^{-4}
β_4	0	0	v	0.023	-10^{-4}
β_5	0	0	h	-0.258	0.011
β_6	0	0			

because it increases the forced convection. However, this correlation is shown to be weak because the range in which the velocity changes ($\pm 20\%$) is very small. The fission cross sections have negligible effect on T_{max} . The delayed neutron fractions and the precursors decay constants have zero derivatives because our reduced model assumes them to be constants at any point.

The multiplication factor is mainly affected by the fission cross sections as expected. The fission cross sections of the two lowest energy groups are the most important. This is because of their higher weight (higher nominal values compared to the fast groups with similar flux magnitudes) in computing k_{eff} . The thermal expansion coefficient is negatively correlated with k_{eff} because by increasing β_{th} , the liquid is mixed more, which in turn causes more precursors to move from regions of higher importance to regions of lower importance near the boundaries. The initial power is negatively correlated with k_{eff} due to the negative temperature feedback coefficient of the system. For the same reason, the heat transfer coefficient is positively correlated with k_{eff} . The lid velocity and viscosity have negligible effect on the multiplication factor.

For the global sensitivities, we computed the first order Sobol indices using quasi Monte Carlo method with Sobol sequence sampling (Sobol, 2001). We selected the size of our sampling matrices to be 10^5 , which generates 2 matrices each of dimension $10^5 \times 27$. The first order Sobol indices were then estimated using the estimators recommended by Saltelli et al. (2010). The computed indices

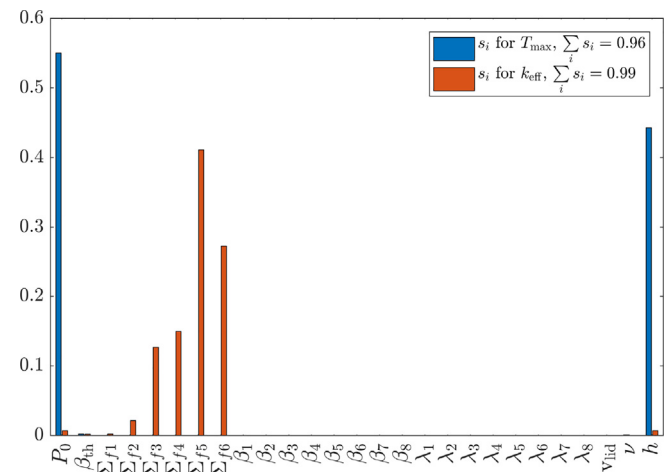


Fig. 10. First order Sobol indices showing the first order sensitivities of T_{max} and k_{eff} to each input parameter. The sum of the first order sensitivities for each output is also shown in the legend.

for both T_{max} and k_{eff} are shown in Fig. 10. The Sobol indices show agreement with the conclusions of the local sensitivities. The maximum temperature is predominantly sensitive to P_0 and h while β_{th} and v have a slight effect on T_{max} . The multiplication factor, on the other hand, is mainly sensitive to the fission cross sections with the

lowest energy groups having the most importance. Although $\Sigma_{f,5}$ has a nominal value of about half $\Sigma_{f,6}$, the Sobol index of $\Sigma_{f,5}$ is about 50% higher than $\Sigma_{f,6}$. This can be explained by the higher flux magnitude of group 5 compared to group 6 as can be seen from the average flux value reported in Table 1 for the nominal case. P_0 and h have a reduced effect while β_{th} has a minimal effect on k_{eff} . The agreement between the local and global sensitivities show that the system is only weakly nonlinear. Additionally, the sum of the computed first order Sobol indices was found to be very close to one, which indicates that second and higher order interactions between the parameters are almost negligible. This confirms the weak nonlinearity of the model.

In total, 3×10^6 model evaluations were performed to complete the uncertainty and sensitivity analysis study. The time to perform these simulations using the reduced model was about 45 min on a personal computer, which is about half the time to perform a single simulation of the full model on the computer cluster. Using 472 snapshots computed in the offline phase, we obtained a gain of about a factor of 5×10^6 in the online computations with respect to the reference model. This demonstrates the advantage of ROM for such applications.

8. Conclusions

The developed ROM algorithm (aPOD) based on POD and the adaptive sparse grids method was applied to a coupled model of a test case for the Molten Salt Fast Reactor. We selected 27 input parameters to model their effect on the distribution of the flux and temperature, and the value of the multiplication factor. In a completely nonintrusive manner, aPOD was able to build a representative (1% accurate) ROM model with 4495 model evaluations. This number was effectively reduced by a factor of 10 with the setting $\mu = 0$. This great reduction was successfully achieved due to the ability of the algorithm to automatically recognize that the 16 dimensions corresponding to β_i and λ_i have no significant effect within the defined range. It was also observed that the initial power, thermal expansion coefficient, fission cross section of the fast 4 groups, lid velocity, and viscosity all have piecewise linear effect on the outputs. On the other hand, the fission cross section for the 2 lowest energy groups and the heat transfer coefficient have slight nonlinear effect. As a test of the model, 1000 Latin Hypercube Sampling points were tested and compared with respect to the reference model. The errors were found to be well within the defined tolerances for all outputs. The multiple-outputs approach was found to require more sampling points to satisfy the desired tolerances compared to a single separate run for each output. This can be explained by the fact that with the single-output ROM model, the POD modes are tailored to that output field and the algorithm only needs to sample points to satisfy the tolerance for that field. The multiple-outputs approach requires the composite POD modes to represent all output fields, which leads to more sampling points to satisfy the tolerances. However, because of the additional sampling in the construction of the reduced model, the error was found to be lower for the multiple approach compared to the single approach.

For an application of uncertainty and sensitivity analysis, we studied the effect of the 27 input parameters on the maximum temperature and the multiplication factor. The density histograms showed a normal distributions of these variables, which can be explained by the uniform distribution assumption of the selected parameters and the weak nonlinearity of the model with respect to the input parameters within the defined ranges. The maximum temperature was shown to be sensitive to the initial power and the heat transfer coefficient while the multiplication factor was mainly sensitive to the fission cross sections as expected. The uncertainty

and sensitivity study was performed using a total of 3 million random points, which were completed in about half the time to run a single simulation of the reference model. The nonintrusive approach of the algorithm provides great potential for studies of complex coupled nuclear systems such as the molten salt reactor, particularly in applications of uncertainty quantification, sensitivity analysis, fuel management, design optimization, and control.

Declaration of Competing Interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgements

F. Alsayyari was supported by King Abdulaziz City for Science and Technology (KACST). M. Tiberger, D. Lathouwers, and J.L. Kloosterman received funding for this project from the Euratom research and training programme 2014–2018 under grant agreement No 661891. The authors would like to thank Dr. M. Aufiero, Dr. P. Rubiolo, and Dr. A. Laureau for providing the necessary data to perform calculations of the molten salt system in the nominal state as described in Section 5.

Appendix A. Supplementary data

Supplementary data associated with this article can be found, in the online version, at <https://doi.org/10.1016/j.anucene.2020.107321>.

References

- Allibert, M., Aufiero, M., Brovchenko, M., Delpech, S., Ghetta, V., Heuer, D., Laureau, A., Merle-Lucotte, E., 2016. Molten salt fast reactors. In: Piro, I.L. (Ed.), Handbook of Generation IV Nuclear Reactors. Woodhead Publishing, pp. 157–188. <https://doi.org/10.1016/B978-0-08-100149-3.00007-0>.
- Alsayyari, F., Perko, Z., Lathouwers, D., Kloosterman, J.L., 2019. A nonintrusive reduced order modelling approach using proper orthogonal decomposition and locally adaptive sparse grids. J. Comput. Phys. 399, <https://doi.org/10.1016/j.jcp.2019.108912>.
- Antoulas, A.C., Sorensen, D.C., Gugercin, S., 2001. A survey of model reduction methods for large-scale systems. Contemp. Math. 280, 193–219.
- Audouze, C., Vuyst, F.D., Nair, P.B., 2009. Reduced-order modeling of parameterized PDEs using time-space-parameter principal component analysis. Int. J. Numer. Meth. Eng. 80, 1025–1057. <https://doi.org/10.1002/nme.2540>.
- Aufiero, M., Rubiolo, P., 2018. Testing and verification of multiphysics tools for fast-spectrum MSRs: the CNRS benchmark. In: Transactions of the 2018 ANS Annual Meeting, Philadelphia, PA, USA.
- Barthelmann, V., Novak, E., Ritter, K., 2000. High dimensional polynomial interpolation on sparse grids. Adv. Comput. Math. 12, 273–288.
- Benner, P., Gugercin, S., Willcox, K., 2015. A survey of projection-based model reduction methods for parametric dynamical systems. SIAM Rev. 57, 483–531.
- Buchan, A.G., Pain, C.C., Fang, F., Navon, I.M., 2013. A POD reduced-order model for eigenvalue problems with application to reactor physics. Int. J. Numer. Meth. Eng. 95, 1011–1032. <https://doi.org/10.1002/nme.4533>.
- Buljak, V., 2011. Inverse Analyses with Model Reduction: Proper Orthogonal Decomposition in Structural Mechanics. Springer, Berlin.
- Bungartz, H.-J., Griebel, M., 2004. Sparse grids. Acta Numer. 13, 147–269. <https://doi.org/10.1017/s0962492904000182>.
- German, P., Ragusa, J.C., 2019. Reduced-order modeling of parameterized multi-group diffusion k-eigenvalue problems. Ann. Nucl. Energy 134, 144–157. <https://doi.org/10.1016/j.anucene.2019.05.049>.
- Griebel, M., 1998. Adaptive sparse grid multilevel methods for elliptic pdes based on finite differences. Computing 61, 151–179.
- Hesthaven, J., Ubbiali, S., 2018. Non-intrusive reduced order modeling of nonlinear problems using neural networks. J. Comput. Phys. 363, 55–78.
- Klimke, A., 2006. Uncertainty Modeling using Fuzzy Arithmetic and Sparse Grids (Ph.D. thesis). Universität Stuttgart, Stuttgart.
- Laureau, A., 2015. Développement de modèles neutroniques pour le couplage thermohydraulique du MSFR et le calcul de paramètres cinétiques effectifs (Ph. D. thesis). Grenoble Alpes University, France. URL: <http://www.theses.fr/2015GREAI064/document>.
- Lorenzi, S., Cammi, A., Luzzi, L., Rozza, G., 2017. A reduced order model for investigating the dynamics of the gen-IV LFR coolant pool. Appl. Math. Model. 46, 263–284. <https://doi.org/10.1016/j.apm.2017.01.066>.

- Ly, H.V., Tran, H.T., 2001. Modeling and control of physical processes using proper orthogonal decomposition. *Math. Comput. Model.* 33, 223–236.
- Manthey, R., Knospe, A., Lange, C., Hennig, D., Hurtado, A., 2019. Reduced order modeling of a natural circulation system by proper orthogonal decomposition. *Prog. Nucl. Energy* 114, 191–200. <https://doi.org/10.1016/j.pnucene.2019.03.010>.
- Nguyen, N., Peraire, J., 2016. Gaussian functional regression for output prediction: model assimilation and experimental design. *J. Comput. Phys.* 309, 52–68.
- Ripley, B.D., 1996. *Pattern Recognition and Neural Networks*. Cambridge University Press. <https://doi.org/10.1017/cbo9780511812651>.
- Saltelli, A., Annoni, P., Azzini, I., Campolongo, F., Ratto, M., Tarantola, S., 2010. Variance based sensitivity analysis of model output. design and estimator for the total sensitivity index. *Comput. Phys. Commun.* 181, 259–270. <https://doi.org/10.1016/j.cpc.2009.09.018>.
- Sartori, A., Baroli, D., Cammi, A., Chiesa, D., Luzzi, L., Ponciroli, R., Previtali, E., Ricotti, M.E., Rozza, G., Sisti, M., 2014. Comparison of a modal method and a proper orthogonal decomposition approach for multi-group time-dependent reactor spatial kinetics. *Ann. Nucl. Energy* 71, 217–229.
- Sobol, I.M., 2001. Global sensitivity indices for nonlinear mathematical models and their monte carlo estimates. *Math. Comput. Simul.* 55, 271–280. [https://doi.org/10.1016/s0378-4754\(00\)00270-6](https://doi.org/10.1016/s0378-4754(00)00270-6).
- Tiberga, M., Lathouwers, D., Kloosterman, J.L., 2019. A discontinuous Galerkin FEM multiphysics solver for the Molten Salt Fast Reactor. In: *International Conference on Mathematics and Computational Methods applied to Nuclear Science and Engineering (M&C 2019)*, Portland, OR, USA.