

Kernel principal component analysis for stochastic input model generation

Xiang Ma, Nicholas Zabaras¹

Materials Process Design and Control Laboratory, Sibley School of Mechanical and Aerospace Engineering, 101 Frank H.T. Rhodes Hall, Cornell University, Ithaca, NY 14853-3801, USA

Abstract

Stochastic analysis of random heterogeneous media provides useful information only if realistic input models of the material property variations are used. These input models are often constructed from a set of experimental samples of the underlying random field. To this end, the Karhunen-Loève (K-L) expansion, also known as principal component analysis (PCA), is the most popular model reduction method due to its uniform mean-square convergence. However, it only projects the samples onto an optimal linear subspace, which results in an unreasonable representation of the original data if they are non-linearly related to each other. In other words, it only preserves the second-order statistics (covariance) of a random field, which is insufficient for reproducing complex structures. This paper applies kernel principal component analysis (KPCA) to construct a reduced-order stochastic input model for the material property variation in heterogeneous media. KPCA can be considered as a nonlinear version of PCA. Through use of kernel functions, KPCA further enables the preservation of high-order statistics of the random field, instead of just two-point statistics as in the standard Karhunen-Loève (K-L) expansion. Thus, this method can model non-Gaussian, non-stationary random fields. In addition, polynomial chaos (PC) expansion is used to represent the random coefficients in KPCA which provides a parametric stochastic input model. Thus, realizations, which are consistent statistically with the experimental data, can be generated in an efficient way. We showcase the methodology by constructing a low-dimensional stochastic input model to represent channelized permeability in porous media.

Key words: Stochastic partial differential equations; Data-driven models; Kernel Principal component analysis; Non-linear model reduction; Flows in random porous media;

¹ Corresponding author: Fax: 607-255-1222, Email: zabaras@cornell.edu, URL: <http://mpdc.mae.cornell.edu/>

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE 17 AUG 2010		2. REPORT TYPE		3. DATES COVERED 00-00-2010 to 00-00-2010	
4. TITLE AND SUBTITLE Kernel principal component analysis for stochastic input model generation				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Cornell University, Materials Process Design and Control Laboratory, 101 Frank H.T. Rhodes Hall, Ithaca, NY, 14853-3801				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT Stochastic analysis of random heterogeneous media provides useful information only if realistic input models of the material property variations are used. These input models are often constructed from a set of experimental samples of the underlying random field. To this end, the Karhunen-Lo'eve (K-L) expansion, also known as principal component analysis (PCA), is the most popular model reduction method due to its uniform mean-square convergence. However, it only projects the samples onto an optimal linear subspace, which results in an unreasonable representation of the original data if they are non-linearly related to each other. In other words, it only preserves the second-order statistics (covariance) of a random field, which is insufficient for reproducing complex structures. This paper applies kernel principal component analysis (KPCA) to construct a reduced-order stochastic input model for the material property variation in heterogeneous media. KPCA can be considered as a nonlinear version of PCA. Through use of kernel functions, KPCA further enables the preservation of high-order statistics of the random field, instead of just two-point statistics as in the standard Karhunen-Lo'eve (K-L) expansion. Thus, this method can model non-Gaussian, non-stationary random fields. In addition, polynomial chaos (PC) expansion is used to represent the random coefficients in KPCA which provides a parametric stochastic input model. Thus, realizations, which are consistent statistically with the experimental data, can be generated in an efficient way. We showcase the methodology by constructing a low-dimensional stochastic input model to represent channelized permeability in porous media.					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 31	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

1 Introduction

Over the past few decades there has been considerable interest among the scientific community in studying physical processes with stochastic inputs. These stochastic input conditions arise from uncertainties in boundary and initial conditions as well as from inherent random material heterogeneities. Material heterogeneities are usually difficult to quantify since it is physically impossible to know the exact property at every point in the domain. In most cases, only a few statistical descriptors of the property variation or only a set of samples can be experimentally determined. This limited information necessitates viewing the property variation as a random field that satisfies certain statistical properties/correlations, which naturally results in describing the physical phenomena using stochastic partial differential equations (SPDEs).

In the past few years, several numerical methods have been developed to solve SPDEs, such as Monte Carlo (MC) method, perturbation approach [1,2], generalized polynomial chaos expansion (gPC) [3–6] and stochastic collocation method [7–13]. However, implicit in all these developments is the assumption that the uncertainties in the SPDEs have been accurately characterized as random variables or random fields through the finite-dimensional noise assumption [11]. The most common choice is the Karhunen-Loève (K-L) expansion [2,3,14,15], which is also known as linear principal component analysis or PCA [16]. Through K-L expansion, the random field can be represented as a linear combination of the deterministic basis functions (eigenfunctions) and the corresponding uncorrelated random variables (random coefficients). The computation of the K-L expansion requires the analytic expression of the covariance matrix of the underlying random field. In addition, the probability distribution of the random variables is always assumed known *a priori*. These two assumptions are obviously not feasible in realistic engineering problems. In most cases, only a few experimentally obtained realizations of the random field are available. Reconstruction techniques are then applied to generate samples of the random field after extracting the statistical properties of the random field through these limited experimental measurements. These processes are quite expensive and numerically demanding if thousands of samples are needed. This leads to the problem of probabilistic model identification or stochastic input model generation, where the purpose is to find a parametric representation of the random field through only limited experimental data.

To this end, a polynomial chaos (PC) representation of the random field through experimental data was first proposed in [17] and improved in subsequent papers [18–23]. This framework consists of three steps: (1) Computing the covariance matrix of the data numerically using the available experimental data; (2) Stochastic reduced-order modeling with a K-L expansion to obtain a set of deterministic basis (eigenfunctions) and the corresponding random

expansion coefficients (called K-L projected random variables); (3) A polynomial chaos representation of these random coefficients is constructed given the realizations of these coefficients which are calculated from data. These realizations are then used for the estimation of the deterministic coefficients in the PC representation, where several methods have been proposed. In the pioneering work [17], maximum likelihood estimation was used to find the PC representation of the K-L random variables. In [18], a Bayesian inference approach was used instead to construct the PC representation of the random field. However, these two works did not take into account the dependencies between various components of the K-L projected random variables. In [19], the Rosenblatt transformation [24] was used to capture these dependencies and maximum entropy approach together with Fisher information matrix was used for the estimation of the PC coefficients. In [20,21,25], a non-intrusive projection approach with Rosenblatt transformation was developed for the PC estimation. Apart from PC representation, in [26], the distribution of the K-L random variables was assumed directly to be uniform within the range of the realizations of these coefficients from data. Later, in [27], the distribution of the K-L random variables was still assumed to be uniform. However, the range of the uniform distribution was found through enforcing the statistical constraints of the random field and solving the resulting optimization problems. In [28], the uncertain input parameters are modeled as independent random variables, whose distributions are estimated using a diffusion-mixed-based estimator. Except the work in [28], all the previous developments rely heavily on the K-L expansion for the reduced-order modeling. However, the K-L expansion has one major drawback. The K-L expansion based stochastic model reduction scheme constructs the *closest linear subspace* of the high-dimensional input space. In other words, it only preserves the covariance of the random field, and therefore is suitable for multi-Gaussian random fields. But most of the random samples contain essential non-linear structures, e.g. higher-order statistics. This directly translates into the fact that the standard linear K-L expansion tends to over-estimate the actual intrinsic dimensionality of the underlying random field. Hence, one needs to go beyond a linear representation of the high-dimensional input space to accurately access the effect of its variability on the output variables.

To resolve this issue, the authors in [29] proposed a nonlinear model reduction strategy for generating data-driven stochastic input models. This method is based on the manifold learning method, where the principal of multidimensional scaling (MDS) is utilized to map the space of viable property variations to a low-dimensional region. Then isometric mapping from this high-dimensional space to a low-dimensional, compact, connected set is constructed via preserving the geodesic distance between data using the IsoMap algorithm [30]. However, this method has two major issues. First, after dimensionality reduction, it only gives us a set of low-dimensional points. It does not give us the inherent patterns (similar to the eigenfunctions as in the K-L

expansion) in the embedded random space. Therefore, it cannot provide us a mathematical parametric input model as in the K-L expansion, i.e. we want to find the form $\mathbf{y} = f(\boldsymbol{\xi})$, where vector $\mathbf{y}$ is a realization of a discrete random field, and vector $\boldsymbol{\xi}$, of dimension much smaller than the original input stochastic dimension, is a set of independent random variables with a specified distribution. In addition, when new experimental data becomes available, this nonlinear mapping needs to be recomputed. Secondly, the IsoMap algorithm requires the computation of the geodesic distance matrix among data. In general, this matrix may be not well defined for real data. Even if it is well defined, the algorithm is computationally expensive.

Both problems of the K-L expansion and the non-linear model reduction algorithm in [29] can be solved with kernel principal component analysis (KPCA), which is a nonlinear form of PCA [31,32]. KPCA has proven to be a powerful tool as a nonlinear feature extractor of classification algorithm [31], pattern recognition [33], image-denoising [34] and statistical shape analysis [35]. The basic idea is to map the data in the input space to a feature space F through some nonlinear map Φ , and then apply the linear PCA there. Through the use of a suitably chosen kernel function, the data becomes more linearly related in the feature space F . In the context of stochastic modeling, there are two pioneered works [36,37]. In [36], KPCA was used to construct the prior model of the unknown permeability field and then gradient-based algorithm is used to solve the history-matching problem. However, the random expansion coefficients of the linear PCA in the feature space are assumed i.i.d. standard normal random variables. This choice clearly does not capture the statical information from the data. In [37], KPCA is used in the feature space F for the selection of a subset of representative realizations containing similar properties to the larger set.

Motivated by all the above mentioned works, in this paper, a stochastic non-linear model reduction framework based on KPCA is developed. To be specific, the KPCA is first used to construct the stochastic reduced-order model in the feature space. The random coefficients of the linear PCA are then represented via PC expansion. Because the K-L expansion is performed in the feature space, the resulting realizations lie in the feature space, and therefore, the mapping from the feature space back to the input space is needed. This is called the “pre-image problem” [34,38]. The pioneering work in solving the pre-image problem is Mika’s fixed-point iterative optimization algorithm [34]. However, this method suffers from numerical instabilities. It is sensitive to the initial guess and is likely to get stuck in local minima. Therefore, it is not suitable for our stochastic modeling. It is noted that this algorithm was also used in the work [36,37] to find the pre-image. In [38], the authors determine a relationship between the distances in the feature space and the distances in the input space. The MDS is used to find the inverse mapping and thus the pre-image. This method is expensive since it involves a singular value

decomposition on a matrix of nearest neighbors. In this work, we propose a new approach to find the pre-image. It is based on local linear interpolation among n -nearest neighbors using only the distances in the input space, which is similar to the method proposed in our earlier work [29].

This paper is organized as follows: In the next section, the mathematical framework of KPCA is considered. In Section 3, the PC representation of the random coefficients is developed. In Section 4, the new pre-image algorithm is introduced. A example with channelized permeability is given in Section 5. Finally, concluding remarks are provided in Section 6.

2 Kernel principal component analysis of random fields

2.1 Problem definition

Let us define a complete probability space $(\Omega, \mathcal{F}, \mathcal{P})$ with sample space Ω which corresponds to the outcomes of some experiments, $\mathcal{F} \subset 2^\Omega$ is the σ -algebra of subsets in Ω and $\mathcal{P} : \mathcal{F} \rightarrow [0, 1]$ is the probability measure. Also, let us define D as a two-dimensional bounded domain. Denote $\mathbf{a}(\mathbf{x}, \omega)$ the random (property) field used to describe and provide a mathematical model of the available experimental data. The random field $\mathbf{a}(\mathbf{x}, \omega)$ in general belongs to an infinite-dimensional probability space. However, in most cases, the random field is associated with a spatial discretization. Thus we can have a finite-dimensional representation of the random field which can be represented/described as a random vector $\mathbf{y} := (y_1, \dots, y_M)^T : \Omega \rightarrow \mathbb{R}^M$. M can be regarded as the number of grid blocks in a discretized model. So each $y_i, i = 1, \dots, M$ is a random variable which represents the random property in each grid block. The dimensionality of the stochastic model is then the length of this vector M . Let $\mathbf{y}_i, i = 1, \dots, N$ be N real column vectors in $\mathbb{R}^M$, i.e. $\mathbf{y}_i \in \mathbb{R}^M$, representing N independent realizations of the random field.

In most cases, M would be a large number. Our problem is to find a reduced-order polynomial chaos representation of this random field that is consistent with the data in some statistical sense. To be specific, we want to find a form $\mathbf{y} = f(\boldsymbol{\xi})$, where vector $\boldsymbol{\xi}$, of dimension much smaller than the original input stochastic dimension M , is a set of independent random variables with a specified distribution. Therefore, by drawing samples $\boldsymbol{\xi}$ in this low-dimensional space, we obtain different realizations of the underlying random field.

2.2 Basic idea of KPCA

Fig. 1 demonstrates the basic idea behind nonlinear kernel PCA. Consider a random field $\mathbf{y} = (y_1, y_2)^T \in \mathbb{R}^2$. If $\mathbf{y}$ is non-Gaussian, y_1 and y_2 can be nonlinearly related to each other in $\mathbb{R}^2$. In this case, linear PCA or K-L expansion attempt to fit a linear surface such that the reconstruction error is minimized (Fig. 1, left). This clearly results in a poor representation of the original data. Now, consider a nonlinear mapping Φ that relates the input space $\mathbb{R}^M$ to another space F

$$\Phi : \mathbb{R}^M \rightarrow F, \quad \mathbf{y} \mapsto \mathbf{Y}. \quad (1)$$

We will refer F as the feature space. In the right figure of Fig. 1, the realizations that are nonlinearly related in $\mathbb{R}^2$ become linearly related in the feature space F . Linear PCA or K-L expansion can now performed in F in order to determine the principal directions in this space.

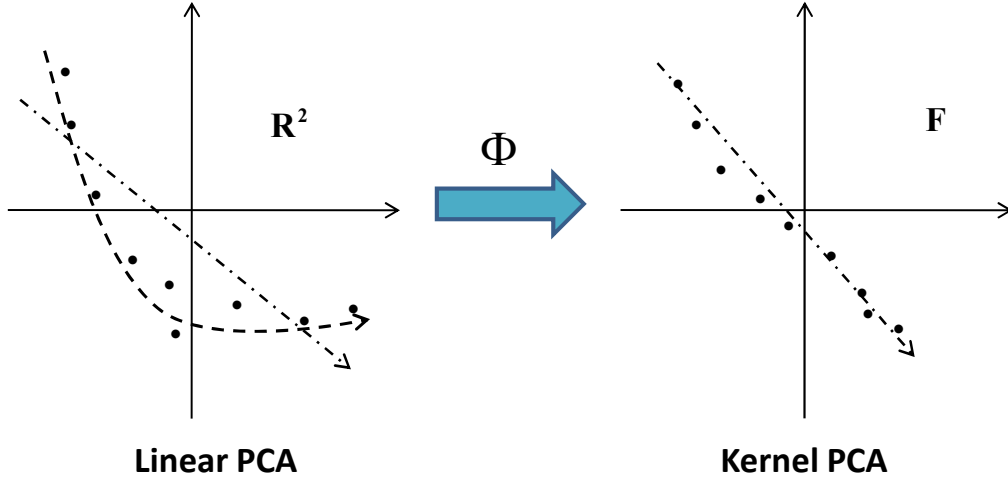


Fig. 1. Basic idea of KPCA. Left: In this non-Gaussian case, the linear PCA is not able to capture the nonlinear relationship among the realizations in the original space. Right: After the nonlinear mapping Φ , the realizations becomes linearly related in the feature space F . Linear PCA or K-L expansion can now be performed in F .

2.3 PCA in feature space

In this section, the theory of KPCA is briefly reviewed. For a detailed introduction, the authors may refer to [31,38].

As shown before, kernel PCA can be considered to be a generalization of linear PCA in the feature space F . Thus, all results on linear PCA can be readily generalized for KPCA. Now, assume we are given N number of realizations

of the random field $\{\mathbf{y}_1, \dots, \mathbf{y}_N\}$, where each realization is represented as a high-dimensional column vector $\mathbf{y}_i \in \mathbb{R}^M$ (e.g. M can be considered as the number of grid blocks in the discretization). The maps of the realizations in the feature space F are $\Phi(\mathbf{y}_i), i = 1, \dots, N$. Denote the mean of the Φ -mapped data by $\bar{\Phi} = \frac{1}{N} \sum_{i=1}^N \Phi(\mathbf{y}_i)$ and define the “centered” map $\tilde{\Phi}$ as

$$\tilde{\Phi} = \Phi(\mathbf{y}) - \bar{\Phi}. \quad (2)$$

Analogous to linear PCA, we need to find eigenvectors $\mathbf{V}$ and eigenvalues λ of the covariance matrix $\mathbf{C}$ in the feature space, where

$$\mathbf{C} = \frac{1}{N} \sum_{i=1}^N \tilde{\Phi}(\mathbf{y}_i) \tilde{\Phi}(\mathbf{y}_i)^T. \quad (3)$$

The dimension of this matrix is $N_F \times N_F$, where N_F is the dimension of the feature space. As explained in [31], N_F could be extremely large. As a result, it will be impossible to compute the $\mathbf{C}$ and solve the eigenvalue problem directly.

Thus, as in [31], a kernel eigenvalue problem is formulated which uses only dot products of vectors in the feature space. We first substitute the covariance matrix into the eigenvalue problem $\mathbf{C}\mathbf{V} = \lambda\mathbf{V}$ and obtain [36]

$$\mathbf{C}\mathbf{V} = \frac{1}{N} \sum_{i=1}^N (\tilde{\Phi}(\mathbf{y}_i) \cdot \mathbf{V}) \tilde{\Phi}(\mathbf{y}_i), \quad (4)$$

which implies that all solutions $\mathbf{V}$ with $\lambda \neq 0$ lie in the span of $\tilde{\Phi}(\mathbf{y}_1), \dots, \tilde{\Phi}(\mathbf{y}_N)$ [31]. Therefore, we can expand the solution $\mathbf{V}$ as [31]

$$\mathbf{V} = \sum_{j=1}^N \alpha_j \tilde{\Phi}(\mathbf{y}_j), \quad (5)$$

and the eigenvalue problem is equivalent to [31]

$$\lambda(\tilde{\Phi}(\mathbf{y}_i) \cdot \mathbf{V}) = (\tilde{\Phi}(\mathbf{y}_i) \cdot \mathbf{C}\mathbf{V}), \quad \forall i = 1, \dots, N. \quad (6)$$

Now, substituting Eq. (5) into Eq. (6), we obtain [31]

$$\lambda \sum_{j=1}^N \alpha_j (\tilde{\Phi}(\mathbf{y}_i) \cdot \tilde{\Phi}(\mathbf{y}_j)) = \frac{1}{N} \sum_{j=1}^N \alpha_j \left(\tilde{\Phi}(\mathbf{y}_i) \cdot \sum_{k=1}^N \tilde{\Phi}(\mathbf{y}_k) \right) (\tilde{\Phi}(\mathbf{y}_k) \cdot \tilde{\Phi}(\mathbf{y}_j)), \quad (7)$$

for $i = 1, \dots, N$. Now define the $N \times N$ kernel matrix $\mathbf{K}$ which is the dot product of vectors in the feature space F :

$$\mathbf{K} : K_{ij} = (\Phi(\mathbf{y}_i) \cdot \Phi(\mathbf{y}_j)). \quad (8)$$

Define the $N \times N$ centering matrix $\mathbf{H} = \mathbf{I} - \frac{1}{N} \mathbf{1}\mathbf{1}^T$, where $\mathbf{I}$ is the $N \times N$ identity matrix and $\mathbf{1} = [1 \dots 1]^T$ is a $N \times 1$ vector. Thus, the centered kernel

matrix $\tilde{\mathbf{K}} : \tilde{K}_{ij} = (\tilde{\Phi}(\mathbf{y}_i) \cdot \tilde{\Phi}(\mathbf{y}_j))$ can be computed as

$$\tilde{\mathbf{K}} = \mathbf{H}\mathbf{K}\mathbf{H}. \quad (9)$$

Substituting Eq. (9) into Eq. (7), we arrive at the following kernel eigenvalue problem [31]:

$$N\lambda\boldsymbol{\alpha} = \tilde{\mathbf{K}}\boldsymbol{\alpha}. \quad (10)$$

In the following, for simplicity, we will denote λ as the eigenvalues of $\tilde{\mathbf{K}}$, i.e. the solutions $N\lambda$ in Eq. (10). We rewrite Eq. (10) in the following matrix form:

$$\boldsymbol{\Lambda}\mathbf{U} = \tilde{\mathbf{K}}\mathbf{U}, \quad (11)$$

where, $\boldsymbol{\Lambda} = \text{diag}(\lambda_1, \dots, \lambda_N)$ is the diagonal matrix of the corresponding eigenvalues and $\mathbf{U} = [\boldsymbol{\alpha}_1, \dots, \boldsymbol{\alpha}_N]$ with $\boldsymbol{\alpha}_i = [\alpha_{i1}, \dots, \alpha_{iN}]^T$ is the matrix containing the eigenvectors.

The eigenvectors need to be normalized, which is $\mathbf{V}_k \cdot \mathbf{V}_k = 1$. Normalizing the solution $\mathbf{V}_i$ in F translates into $\lambda_i(\boldsymbol{\alpha}_i \cdot \boldsymbol{\alpha}_i) = 1$. From the solution of the eigenvalue problem, we can write $\boldsymbol{\alpha}_i \cdot \boldsymbol{\alpha}_i = 1$, since the eigenvectors from the eigenvalue problem in Eq. (11) are normalized. Therefore, if we divide $\boldsymbol{\alpha}_i$ by $\sqrt{\lambda_i}$, we obtain $\boldsymbol{\alpha}_i/\sqrt{\lambda_i} \cdot \boldsymbol{\alpha}_i/\sqrt{\lambda_i} = 1/\lambda_i$.

Therefore, through Eq. (5), the i th orthonormal eigenvector of the covariance matrix $\mathbf{C}$ in the feature space can be shown to be [31,38]

$$\mathbf{V}_i = \sum_{j=1}^N \frac{\alpha_{ij}}{\sqrt{\lambda_i}} \tilde{\Phi}(\mathbf{y}_j) = \sum_{j=1}^N \tilde{\alpha}_{ij} \tilde{\Phi}(\mathbf{y}_j), \text{ where } \tilde{\alpha}_{ij} = \frac{\alpha_{ij}}{\sqrt{\lambda_i}}. \quad (12)$$

Let $\mathbf{y}$ be a realization of the random field, with a mapping $\tilde{\Phi}(\mathbf{y})$ in F . Then $\tilde{\Phi}(\mathbf{y})$ can be decomposed in the following way:

$$\tilde{\Phi}(\mathbf{y}) = \sum_{i=1}^N z_i \mathbf{V}_i + \bar{\Phi}, \quad (13)$$

where z_i is the projection coefficient onto the i th eigenvector $\mathbf{V}_i$:

$$z_i = \mathbf{V}_i \cdot \tilde{\Phi}(\mathbf{y}) = \sum_{j=1}^N \tilde{\alpha}_{ij} (\tilde{\Phi}(\mathbf{y}) \cdot \tilde{\Phi}(\mathbf{y}_j)). \quad (14)$$

2.4 Computing dot products in feature space

From Eq. (9), it is seen that in order to compute the kernel matrix, only the dot products of vectors in the feature space F are required, while the explicit

calculation of the map $\Phi(\mathbf{y})$ does not need to be known. As shown in [31], the dot product can be computed through the use of the kernel function. This is the so called “kernel trick”. Not all arbitrary functions but the Mercer kernels can be used as a kernel function [31]. The kernel function $k(\mathbf{y}_i, \mathbf{y}_j)$ calculates the dot product in space F directly from the vectors of the input space $\mathbb{R}^M$:

$$k(\mathbf{y}_i, \mathbf{y}_j) = (\Phi(\mathbf{y}_i) \cdot \Phi(\mathbf{y}_j)). \quad (15)$$

The commonly used kernel functions are polynomial kernel and Gaussian kernels.

2.4.1 Kernel for linear PCA

If the kernel function is chosen as polynomial kernel of order one

$$k(\mathbf{y}_i, \mathbf{y}_j) = (\mathbf{y}_i \cdot \mathbf{y}_j), \quad (16)$$

then the linear PCA is actually performed on the sample realizations. It is noted that the use of the kernel matrix to perform the linear PCA is actually the same as the “method of snapshots” which is well known in the area of reduced order modeling [39]. This method is more computationally efficient than the standard implementation of the K-L expansion as in [17]. Using the kernel matrix, only a eigenvalue problem of size $N \times N$ is needed, whereas in [17] the size of the eigenvalue problem is $M \times M$. In most cases, the number of available experimental data is much smaller than the dimensionality of the data itself.

2.4.2 Kernel for nonlinear PCA

Choosing a nonlinear kernel function leads to performing nonlinear PCA. The most common kernel function is the Gaussian kernel:

$$k(\mathbf{y}_i, \mathbf{y}_j) = \exp\left(-\frac{\|\mathbf{y}_i - \mathbf{y}_j\|^2}{2\sigma^2}\right), \quad (17)$$

where $\|\mathbf{y}_i - \mathbf{y}_j\|^2$ is the squared L_2 -distance between two realizations. The kernel width parameter σ controls the flexibility of the kernel. A larger value of σ allows more “mixing” between elements of the realizations, whereas a smaller value of σ uses only a few significant realizations. As recommended in [35], a typical choice for σ is the average minimum distance between two realizations in the input space:

$$\sigma^2 = c \frac{1}{N} \sum_{i=1}^N \min_{j \neq i} \|\mathbf{y}_i - \mathbf{y}_j\|^2, \quad j = 1, \dots, N, \quad (18)$$

where c is a user-controlled parameter.

2.5 Stochastic reduced-order modeling via KPCA

Since only the dot products of vectors in the feature space are available through use of the kernel function, we first rewrite Eq. (14) in terms of the kernel function [38]:

$$z_i = \sum_{j=1}^N \tilde{\alpha}_{ij} \left(\tilde{\Phi}(\mathbf{y}) \cdot \tilde{\Phi}(\mathbf{y}_j) \right) = \sum_{j=1}^N \tilde{\alpha}_{ij} \tilde{k}(\mathbf{y}, \mathbf{y}_j), \quad (19)$$

where

$$\tilde{k}(\mathbf{y}, \mathbf{y}_j) = k(\mathbf{y}, \mathbf{y}_j) - \frac{1}{N} \mathbf{1}^T \mathbf{k}_y - \frac{1}{N} \mathbf{1}^T \mathbf{k}_{y_j} + \frac{1}{N^2} \mathbf{1}^T \mathbf{K} \mathbf{1}, \quad (20)$$

with

$$\mathbf{k}_y = [k(\mathbf{y}, \mathbf{y}_1), \dots, k(\mathbf{y}, \mathbf{y}_N)]^T. \quad (21)$$

We can write Eq. (19) in a matrix form:

$$\mathbf{Z} = \mathbf{A}^T \mathbf{k}_y + \mathbf{b}, \quad (22)$$

where $\mathbf{A} = \mathbf{H} \tilde{\boldsymbol{\alpha}}$ and $\mathbf{b} = -\frac{1}{N} \tilde{\boldsymbol{\alpha}}^T \mathbf{H} \mathbf{K} \mathbf{1}$.

Now suppose the eigenvectors are ordered by decreasing eigenvalues and we can only work in the low-dimensional subspace which is spanned by the first r largest eigenvectors, where, in general, $r \ll N$. Then the decomposition in Eq. (13) can be truncated after the first r terms:

$$\tilde{\Phi}_r(\mathbf{y}) \approx \sum_{i=1}^r z_i \mathbf{V}_i + \bar{\Phi} = \sum_{i=1}^N \beta_i \Phi(\mathbf{y}_i) \quad (23)$$

where β_i is the i th component of the vector $\boldsymbol{\beta} = \mathbf{A} \mathbf{Z} + \frac{1}{N} \mathbf{1}$. It is noted that since only the first r eigenvectors are used, $\tilde{\boldsymbol{\alpha}}$ used in Eq. (22) only contains the first r columns of the original matrix.

Then the stochastic reduced-order input model in the feature space can be defined as: for any realization $\mathbf{Y} \in F$, we have

$$\mathbf{Y}_r = \sum_{i=1}^N \beta_i \Phi(\mathbf{y}_i), \quad \text{with } \boldsymbol{\beta} = \mathbf{A} \boldsymbol{\xi} + \frac{1}{N} \mathbf{1}. \quad (24)$$

Here, the subscript r is to emphasize that the realization $\mathbf{Y}_r$ is reconstructed using only the first r eigenvectors. $\boldsymbol{\xi} := [\xi_1, \dots, \xi_r]^T$ is a r -dimensional random vector.

According to the properties of the K-L expansion [3,15], the random vectors $\boldsymbol{\xi}$ satisfy the following two conditions:

$$E[\xi_i] = 0, \quad E[\xi_i \xi_j] = \delta_{ij} \frac{\lambda_i}{N}, \quad i, j = 1, \dots, r. \quad (25)$$

Therefore, the random coefficients ξ_i are uncorrelated but not independent. The realizations of these random coefficients can be obtained through Eq. (22)

$$\boldsymbol{\xi}^{(i)} = \mathbf{A}^T \mathbf{k}_{\mathbf{y}_i} + \mathbf{b}, \quad i = 1, \dots, N. \quad (26)$$

Our problem then reduces to identify the random vector $\boldsymbol{\xi} := [\xi_1, \dots, \xi_r]^T$, given its N samples $\boldsymbol{\xi}^{(i)} = [\xi_1^{(i)}, \dots, \xi_r^{(i)}]$, $i = 1, \dots, N$. A polynomial chaos representation is introduced in the next section.

Finally, similarly to the K-L expansion [3,15], the minimum mean squared error (MSE) of truncated expansion Eq. (24) in the feature space can be shown to be

$$E(\|\mathbf{Y} - \mathbf{Y}_r\|^2) = \frac{1}{N} \sum_{i=r+1}^N \lambda_i \quad (27)$$

Remark 1. For linear PCA, we only need to replace the kernel function in Eq. (22) with the kernel function for linear PCA Eq. (16) and replace $\Phi(\mathbf{y}_i)$ in Eq. (24) with the data $\mathbf{y}_i$ directly.

3 Polynomial chaos representation of the stochastic reduced order model

The problem is now to identify a random vector $\boldsymbol{\xi} : \Omega \rightarrow \mathbb{R}^r$, given a set of independent samples $\{\boldsymbol{\xi}^{(i)}\}_{i=1}^N$. For this purpose, we use a polynomial chaos (PC) expansion to represent $\boldsymbol{\xi}$. Several methods have been proposed to solve the expansion coefficients in the resulted PC expansion, such as maximum likelihood [17], Bayesian inference [18], maximum entropy method [19] and non-intrusive projection method [21,25]. As mentioned before, the components of random vector $\boldsymbol{\xi}$ are uncorrelated but not independent. Although Rosenblatt transformation can be used to reduce the problem to a set of independent random variables [20,25], it is computationally expensive for high-dimensional problems. In this work, to reduce the computational cost, we further assume the independence between components of $\boldsymbol{\xi}$. This is generally not the case for arbitrary stochastic processes. However, in the works [18] and [21], it has been numerically verified that this strong hypothesis gives accurate results.

3.1 PC expansion of random variables

The theory and properties of the PC expansion have been well documented in various references [3–5]. In this approach, any random variable with finite variance can be expanded in terms of orthogonal polynomials of specific

standard random variables. Since each ξ_i is independent, it can be separately decomposed onto an one-dimensional PC basis of degree p :

$$\xi_i = \sum_{j=0}^p \gamma_{ij} \Psi_j(\eta_i), \quad i = 1, \dots, r, \quad (28)$$

where the η_i are i.i.d. random variables. The random basis functions $\{\Psi_j\}$ are chosen according to the type of random variable $\{\eta_i\}$ that has been used to describe the random input. For example, if Gaussian random variables are chosen then the Askey based orthogonal polynomials $\{\Psi_j\}$ are chosen to be Hermite polynomials, if η_i are chosen to be uniform random variables, then $\{\Psi_j\}$ must be Legendre polynomials [4]. Although the Hermite polynomials are used in this paper, the method developed can be applied to generalized polynomial chaos expansions. The Hermite polynomials are given by

$$\begin{aligned} \Psi_0(\eta_i) &= 1, \quad \Psi_1(\eta_i) = \eta_i, \\ \Psi_{j+1}(\eta_i) &= \eta_i \Psi_j(\eta_i) - j \Psi_{j-1}(\eta_i), \quad \text{if } j \geq 1. \end{aligned} \quad (29)$$

The above one-dimensional Hermite polynomials are orthogonal with respect the corresponding probability density function (PDF) of the standard normal random variable:

$$E[\Psi_i \Psi_j] = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{+\infty} \Psi_i(\eta) \Psi_j(\eta) e^{-\frac{\eta^2}{2}} d\eta = i! \delta_{ij}. \quad (30)$$

Thus, the PC coefficients can be computed through Galerkin projection:

$$\gamma_{ij} = \frac{E[\xi_i \Psi_j(\eta)]}{E[\Psi_j^2(\eta)]} = \frac{1}{\sqrt{2\pi} j!} \int_{-\infty}^{+\infty} \xi_i \Psi_j(\eta) e^{-\frac{\eta^2}{2}} d\eta, \quad i = 1, \dots, r, \quad j = 0, \dots, p. \quad (31)$$

A numerical integration is needed to evaluate this integral. However, it is noted that the random variable ξ does not belong to the same stochastic space as η , a non-linear mapping $\Gamma : \eta \rightarrow \xi$ is thus needed which preserves the probabilities such that $\Gamma(\eta)$ and ξ have the same distributions. Here, a non-intrusive projection method using empirical cumulative distribution functions (CDFs) of samples, which was developed in [21], is utilized to compute the PC coefficients.

3.2 A non-intrusive projection method for calculating PC coefficients

The non-linear mapping $\Gamma : \eta \rightarrow \xi$ can be defined by employing the Rosenblatt transformation [24] as shown below for each ξ_i :

$$\xi_i \stackrel{d}{=} \Gamma_i(\eta_i), \quad \Gamma_i \equiv F_{\xi_i}^{-1} \circ F_{\eta_i}, \quad (32)$$

where F_{ξ_i} and F_{η_i} denote the CDFs of ξ_i and η_i , respectively. Here, the equalities, “ $\stackrel{d}{=}$ ” should be interpreted in the sense of distribution such that the PDFs of random variables on both sides are equal. The statistical toolbox of MATLAB provides functions to evaluate the CDF of many standard PC random variables.

However, the marginal CDF of the samples ξ_i is not known and can only be evaluated numerically from the available data. Here, the kernel density estimation approach is used to construct the empirical CDF of ξ_i [40]. Now, let $\{\xi_i^{(s)}\}_{s=1}^N$ be N samples of ξ_i obtained from Eq. (26). Then the marginal pdf of ξ_i is evaluated as:

$$p_{\xi_i}(\xi) \approx \frac{1}{N} \sum_{s=1}^N \frac{1}{\sqrt{2\pi}\tau} \exp\left(-\frac{\xi - \xi_i^{(s)}}{2\tau^2}\right), \quad (33)$$

where the bandwidth τ should be chosen to balance smoothness and accuracy. Then CDF of ξ can be obtained by integrating Eq. (33) and then the inverse CDF is computed. The MATLAB function, `ksdensity`, in statistical tool box is used to find the inverse CDF of ξ_i , where the τ is automatically computed from the information of data.

Now the expectation in the Galerkin projection formulate Eq. (31) is well defined and can be computed in η -space. Then, the coefficients of the PC expansion can be computed using a Gauss-Hermite quadrature:

$$\gamma_{ij} = \frac{1}{\sqrt{2\pi}j!} \int_{-\infty}^{+\infty} \xi_i \Psi_j(\eta) e^{-\frac{\eta^2}{2}} d\eta \approx \frac{1}{\sqrt{\pi}j!} \sum_{k=1}^{N_g} \omega_k \Gamma_i(\sqrt{2}\mu_k) \Psi_j(\sqrt{2}\mu_k), \quad (34)$$

where the $\{\omega_k, \mu_k\}_{k=1}^{N_g}$ are integration weights and points of Gauss-Hermite quadrature. It is noted here that a transformation $\eta = \sqrt{2}\mu$ is used since the weight in the Gauss-Hermite quadrature is $e^{-\eta^2}$ while the PDF of Gauss random variable is $\frac{1}{\sqrt{2\pi}}e^{-\eta^2/2}$.

3.3 PC representation of the random field

Now, putting both KPCA and PC expansion together, one arrives at the following r -dimensional PC representation of the stochastic random field in the feature space:

$$\mathbf{Y}_r = \sum_{i=1}^N \beta_i \Phi(\mathbf{y}_i), \quad \text{with } \boldsymbol{\beta} = \mathbf{A}\boldsymbol{\zeta} + \frac{1}{N}\mathbf{1}, \quad (35)$$

where the $r \times 1$ column vector $\boldsymbol{\zeta}$ is:

$$\boldsymbol{\zeta} = \left[\sum_{j=0}^p \gamma_{1j} \Psi_j(\eta_1), \dots, \sum_{j=0}^p \gamma_{rj} \Psi_j(\eta_r) \right]^T. \quad (36)$$

Therefore, by drawing i.i.d. samples of standard Gaussian random variables $\eta_i, i = 1, \dots, r$, we obtain different realizations of the underlying random field in the feature space F . Now the dimensionality of the stochastic input space is successfully reduced from a large number M to a small number r .

4 The pre-image problem in KPCA

The simulated realizations of the random field from Eq. (35) are in the feature space F . However, we are interested in obtaining realizations in the physical input space. Therefore, an inverse mapping is required as $\mathbf{y} = \Phi^{-1}(\mathbf{Y})$. This is known as the pre-image problem [34,38]. As demonstrated in [34], this pre-image may not exist or if it exists, it may be not unique. Therefore, we can only settle for an approximate pre-image $\hat{\mathbf{y}}$ such that $\Phi(\hat{\mathbf{y}}) \approx \mathbf{Y}_r$.

4.1 Fixed-point iteration for finding the pre-image

One solution to the pre-image problem is to address this problem as a nonlinear optimization problem by minimizing the squared distance in the feature space between $\Phi(\hat{\mathbf{y}})$ and $\mathbf{Y}_r$:

$$\rho(\hat{\mathbf{y}}) = \|\Phi(\hat{\mathbf{y}}) - \mathbf{Y}_r\|^2. \quad (37)$$

The extremum can be obtained by setting $\nabla_{\hat{\mathbf{y}}} \rho = 0$. For Gaussian kernel Eq. (17), this nonlinear optimization problem can be solved by a fixed-point iteration method [34]:

$$\hat{\mathbf{y}}_{t+1} = \frac{\sum_{i=1}^N \beta_i \exp(-\|\hat{\mathbf{y}}_t - \mathbf{y}_i\|^2/(2\sigma^2)) \mathbf{y}_i}{\sum_{i=1}^N \beta_i \exp(-\|\hat{\mathbf{y}}_t - \mathbf{y}_i\|^2/(2\sigma^2))}. \quad (38)$$

As can be easily seen, the pre-image in this case will depend on the initial starting point and is likely to get stuck in local minima. In addition, this scheme is numerically unstable and one has to try a number of initial guesses. Therefore, it is not suitable for our stochastic simulation since we need to have a one to one mapping and to find an efficient way for generating a large number of samples.

Furthermore, it is noted that the pre-image obtained in this scheme is in the span of all realizations $\mathbf{y}_i$'s, i.e. it is a linear combination of all the available

realizations:

$$\hat{\mathbf{y}} = \sum_{i=1}^N \theta_i \mathbf{y}_i, \quad \text{with} \quad \sum_{i=1}^N \theta_i = 1, \quad (39)$$

where the weights

$$\theta_i = \frac{\beta_i \exp(-\|\hat{\mathbf{y}}_t - \mathbf{y}_i\|^2/(2\sigma^2))}{\sum_{i=1}^N \beta_i \exp(-\|\hat{\mathbf{y}}_t - \mathbf{y}_i\|^2/(2\sigma^2))}. \quad (40)$$

Due to the exponential term in the Gaussian kernel, the contributions of the realizations typically drop rapidly with increasing distance from the pre-image. In other words, the influence of training realizations with smaller distance to $\hat{\mathbf{y}}$ will tend to be bigger. Therefore, it is reasonable to use only nearest neighbors of the pre-image for local linear interpolation.

4.2 Local linear interpolation for finding the pre-image

As illustrated in the last section, we can find the pre-image using local linear interpolation within n -nearest neighbors. Motivated by our previous work in [29], the Euclidean distances are used as the interpolation weights. Actually, there exists a simple relationship between the feature-space and input-space distance for Gaussian kernel [38,41]. The basic idea of the method is shown in Fig. 2. For an arbitrary realization $\mathbf{Y}_r$ in F , we first calculate its distances to the nearest neighbors in the feature space. Then the distances in the input space are recovered. Finally, the input-space distances are used as the local linear interpolation weights.

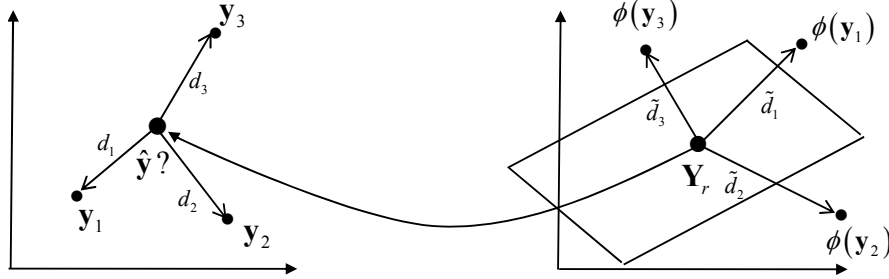


Fig. 2. Basic idea of the proposed pre-image method.

Given any realization $\mathbf{Y}_r \in F$, we can compute the squared feature distance between $\mathbf{Y}_r$ to the i th mapped data as:

$$\tilde{d}_i^2(\mathbf{Y}_r, \Phi(\mathbf{y}_i)) = \|\mathbf{Y}_r - \Phi(\mathbf{y}_i)\|^2 = \|\Phi(\mathbf{y}_i)\|^2 + \|\mathbf{Y}_r\|^2 - 2\mathbf{Y}_r^T \Phi(\mathbf{y}_i), \quad (41)$$

for $i = 1, \dots, N$. Recall that for Gaussian kernel, $k(\mathbf{y}_i, \mathbf{y}_i) = 1$ and $\mathbf{Y}_r = \sum_{i=1}^N \beta_i \Phi(\mathbf{y}_i)$. Then for each feature distance $\tilde{d}_i^2(\mathbf{Y}_r, \Phi(\mathbf{y}_i))$, $i = 1, \dots, N$, we

can compute them in the following matrix form:

$$\tilde{\mathbf{d}}^2 = \mathbf{1} + \boldsymbol{\beta}^T \mathbf{K} \boldsymbol{\beta} - 2\mathbf{K}\boldsymbol{\beta}, \quad (42)$$

where the vector $\tilde{\mathbf{d}}^2 = [\tilde{d}_1^2, \dots, \tilde{d}_N^2]^T$. We thus can sort this vector in ascending order to identify the n -nearest neighbors with respect to $\mathbf{Y}_r$, $\Phi(\tilde{\mathbf{y}}_i), i = 1, \dots, n$.

On the other hand, the squared feature distance between the Φ -map of the pre-image $\hat{\mathbf{y}}$ and the i th mapped data is:

$$\begin{aligned} \hat{d}_i^2(\Phi(\hat{\mathbf{y}}), \Phi(\mathbf{y}_i)) &= \|\Phi(\hat{\mathbf{y}}) - \Phi(\mathbf{y}_i)\|^2 \\ &= k(\hat{\mathbf{y}}, \hat{\mathbf{y}}) + k(\mathbf{y}_i, \mathbf{y}_i) - 2k(\hat{\mathbf{y}}, \mathbf{y}_i) \\ &= 2(1 - k(\hat{\mathbf{y}}, \mathbf{y}_i)), \end{aligned} \quad (43)$$

for $i = 1, \dots, N$ and where we have used $k(\hat{\mathbf{y}}, \hat{\mathbf{y}}) = k(\mathbf{y}_i, \mathbf{y}_i) = 1$ again since Gaussian kernel is used. Furthermore, we have the squared input-space distance:

$$k(\hat{\mathbf{y}}, \mathbf{y}_i) = \exp\left(-\frac{\|\hat{\mathbf{y}} - \mathbf{y}_i\|^2}{2\sigma^2}\right),$$

from which we obtain

$$d_i^2 = \|\hat{\mathbf{y}} - \mathbf{y}_i\|^2 = -2\sigma^2 \log(k(\hat{\mathbf{y}}, \mathbf{y}_i)), \quad (44)$$

for $i = 1, \dots, N$. Substituting Eq. (43) into Eq. (44), one arrives at

$$d_i^2 = \|\hat{\mathbf{y}} - \mathbf{y}_i\|^2 = -2\sigma^2 \log(1 - 0.5\hat{d}_i^2), \quad (45)$$

for $i = 1, \dots, N$. Because we try to find an approximate pre-image such that $\Phi(\hat{\mathbf{y}}) \approx \mathbf{Y}_r$, it is straightforwardly to identify the relationship $\hat{d}_i^2 \approx \tilde{d}_i^2$ from Eqs. (41) and (43). Therefore, the squared input-distance between the approximate pre-image $\hat{\mathbf{y}}$ and i th realization can be computed by:

$$d_i^2 = \|\hat{\mathbf{y}} - \mathbf{y}_i\|^2 = -2\sigma^2 \log(1 - 0.5\tilde{d}_i^2), \quad (46)$$

for $i = 1, \dots, N$ and where $\tilde{d}_i^2$ is given by Eq. (42).

Finally, the pre-image $\hat{\mathbf{y}}$ for a feature space realization $\mathbf{Y}_r$ is given by

$$\hat{\mathbf{y}} = \frac{\sum_{i=1}^n \frac{1}{\tilde{d}_i} \tilde{\mathbf{y}}_i}{\sum_{i=1}^n \frac{1}{\tilde{d}_i}}, \quad (47)$$

where $\tilde{\mathbf{y}}_i, i = 1, \dots, n$ are the n -nearest neighbors. It is noted that here we use the n -nearest neighbors in the feature space. However, they are the same as the n -nearest neighbors in the input space since Eq. (46) is monotonically increasing.

Therefore, the pre-image $\hat{\mathbf{y}}$ of an arbitrary realization in the feature space is the weighted sum of the pre-images of the n -nearest neighbors of $\mathbf{Y}_r$ in the feature space, where the nearest neighbors are taken from the samples $\mathbf{y}_i, i = 1, \dots, N$. This local linear interpolation procedure is based on the principle that a small region in a highly curved manifold can be well approximated as a linear patch. It is easily verified that the interpolation weights satisfies Eq. (39). Thus, a unique pre-image can now be obtained using simple algebraic calculations in a single step (no iteration is required) and is suitable for stochastic simulation.

5 Numerical example

In this section, we apply Kernel PCA on modeling random permeability field of complex geological channelized structures. This structure is characterized by multipoint statistics (MPS), which expresses the joint variability at many more than two locations [42]. Therefore, only mean and correlation structure (two-point statistics) are not enough to capture the underlying uncertainty of the random field and thus linear PCA or standard K-L expansion is expected to fail.

5.1 Generation of experimental samples

In [17], the experimental data was obtained through solution of stochastic inverse problems given several realizations of the system outputs. This method is computationally expensive if a large number of samples is required. To this end, we utilize the reconstruction technique from an available training image to numerically generate samples of the random field [29]. The basic idea is that the training image contains geological structure or continuity information of the underlying random field, and the MPS algorithm creates realizations of the random field honoring this structural relationship [36]. In this example, one of the MPS algorithms, the single normal equation simulation ‘snesim’ algorithm is used to generate the channelized permeability [42] from the training image shown in Fig. 3. It is a binary image where 1 designates channel sand and 0 designates background mud. Since we are only interested in the geological structure not the value itself, a further assumption is made such that the image is a logarithmic transformation of the original permeability field and the values 1 and 0 are the permeability values themselves.

Using the ‘snesim’ algorithm, 1000 realizations of a channelized permeability field, of dimension 45×45 , are created from the training image in Fig. 3. These serve as the training samples (experimental data) of the random field for the construction of KPCA. Additional set of realizations, which is not included

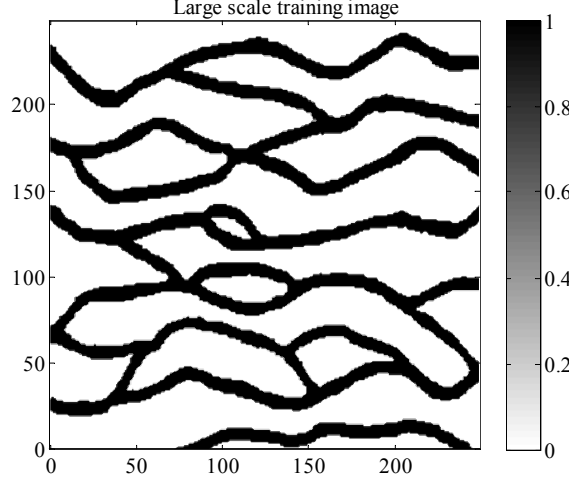


Fig. 3. The large-scale training image for the ‘snesim’ algorithm.

in the training samples, will serve as the test samples to verify the accuracy of the constructed reduced-order model using KPCA algorithm. Some of the realizations created are shown in Fig. 4.

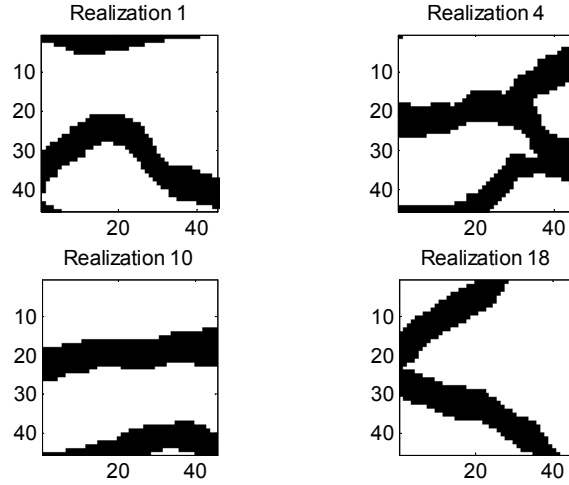


Fig. 4. Some of the realizations created using the ‘snesim’ algorithm.

Each realization of the random field can be considered as a 2025-dimensional vector. Since each vector consists of only 0 and 1, all the samples occupy only corners of a 2025-dimensional hypercube. Therefore, they are not linearly related and linear PCA will fail. This problem is similar to the image-denoising problem in machine learning community where binary images of digits are used and the performance of KPCA is proved to be superior to that of linear PCA [34,38].

5.2 Comparison between linear PCA and Kernel PCA

The kernels Eqs. (16) and (17) are used now to perform linear PCA and Kernel PCA, respectively, on the 1000 sample realizations. The parameter c in Eq. (18) is taken to be 10. The kernel matrices are formulated and subsequent eigenvalue problems are solved. The corresponding first 100 eigenvalues and reconstruction mean squared error (MSE) are shown in Fig. 5. For non-linear data set, the failure of linear PCA means overestimating the intrinsic dimensionality of the data. Therefore, in order to compare the performance from both methods, we need to keep the same number of eigenvectors. A rule of thumb is to choose r such that $\sum_{i=1}^r \lambda_i / \sum_{i=1}^N \lambda_i$ is sufficiently close to one. However, this rule is not applied here. Since if the reconstruction MSE is sufficiently small, there is no need to keep so many eigenvectors. To this end, only the largest 30 eigenvalues are retained for the KPCA, which corresponds to about 75% energy of the random field. The associated reconstruction MSE is 0.003, which is small enough to ensure an accurate expansion in the feature space. On the same time, we also keep 30 eigenvalues for linear PCA, where the reconstruction MSE is 107.242, which is much larger than that of KPCA. Many more eigenvalues are needed to reduce the MSE. This observation partially indicates the failure of linear PCA.

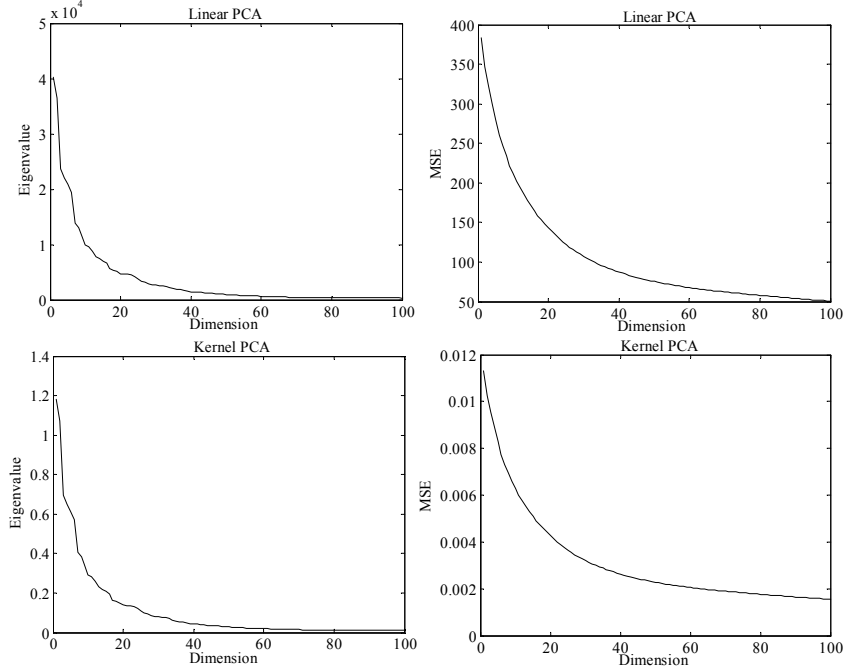


Fig. 5. Plots of the eigenspectrum (left column) and MSE (right column) from linear PCA (top row) and Kernel PCA (bottom row).

We next try to reconstruct one of the test samples using both methods. The results are shown in Fig. 6 with different number of eigenvectors retained. By reconstruction, we mean we first project the test samples onto the

low-dimensional space and compute their coordinates. Then using only the low-dimensional coordinates, the original sample in the input space is reconstructed from linear PCA and Kernel PCA pre-image algorithm, respectively. 10 nearest neighbors are used in finding the pre-image. It is seen that using linear PCA, the reconstruction results improve slowly with increasing number of eigenvectors. A large number of eigenvectors are needed to obtain a good result. This is consistent with our previous discussion that the linear PCA will overestimate the actual dimensionality of the data in the nonlinear case. However, the reconstructed permeability value is still not correct even with $r = 500$. On the other hand, only 30 eigenvectors are needed to get a good reconstruction from KPCA. Increasing the number of eigenvectors will not improve the results since the nearest neighbors have been correctly identified and the reconstruction MSE in the feature space is quite small. The first 9 identified nearest neighbors of the test sample are shown in Fig. 7 with $r = 30$. It is clearly seen that the first 3 neighbors have the geological structures which are most similar to our test sample. This verifies that the introduced pre-image algorithm indeed finds the correct nearest neighbors among the data using only the distances in the features space and input space.

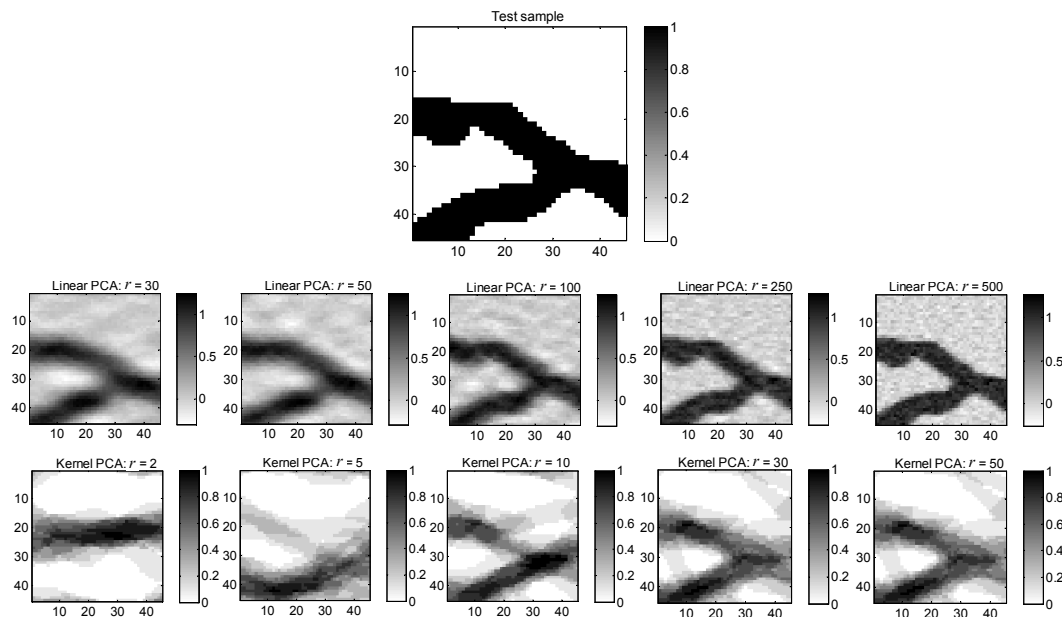


Fig. 6. Reconstruction of a test sample (top figure) using linear PCA (middle row) and Kernel PCA (bottom row) with different number of eigenvectors retained.

To further verify the Kernel PCA algorithm, more test samples are reconstructed. To compare its performance with that of linear PCA, 30 eigenvectors are retained in both cases since it was shown before that $r = 30$ is enough for Kernel PCA. The results are shown in Fig. 8. Again, Kernel PCA consistently gives more accurate results and the reconstruction samples are more like the original binary images.

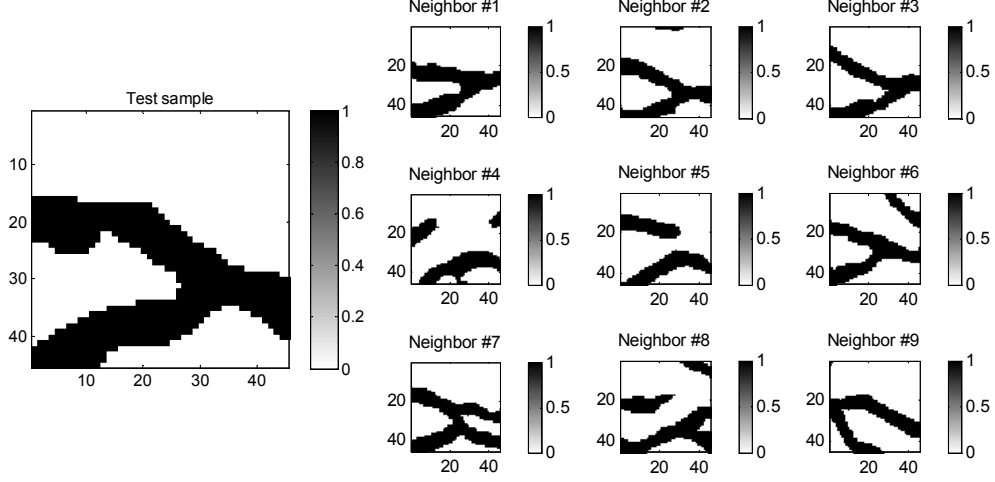


Fig. 7. The first 9 nearest neighbors of the test samples identified in the feature space with $r = 30$.

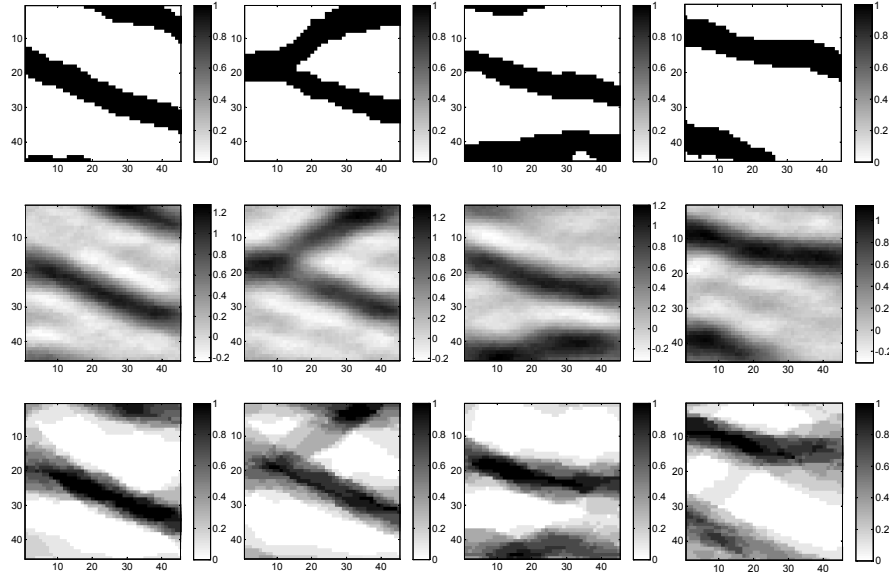


Fig. 8. More reconstruction results of different test samples (top row) using linear PCA (middle row) and Kernel PCA (bottom row) with $r = 30$.

In this section, we compared the performance of linear PCA and Kernel PCA for reconstruction. Although linear PCA is not optimal, it still more or less provided us the desired geological structure. However, what we are interested in is the generation of stochastic realizations. In the next section, we will show that the linear PCA cannot give us arbitrary realizations with expected geological patterns.

5.3 PC representation and stochastic sampling in the low-dimensional space

In this section, we construct the PC representation of the random field using the method introduced in Section 3.2. We keep 30 eigenvectors for both linear PCA and kernel PCA. First, the samples $\xi^{(i)}$ of ξ are computed by inserting the 1000 sample realizations $\mathbf{y}_i$ into Eq. (26). Then these samples are used to construct the inverse marginal CDF for each dimension through kernel density estimation. Finally, Eq. (34) is used to calculate the PC coefficients for each dimension. A PC representation of the random channelized permeability is constructed next, which is only of 30 dimensions, compared with the original 2025-dimensional space. By drawing 30 i.i.d. standard normal random variables η_i from this low-dimensional space and substituting them into Eqs. (35) and (36), any realization of the underlying random permeability field can now be obtained in an inexpensive way.

Fig. 9 depicts 4 different marginal PDFs of the initial and identified random variables using the non-intrusive projection method for Hermite chaos with increasing expansion orders using linear PCA. The marginal PDF of initial random variable is obtained by plotting the kernel density estimation of the 1000 sample coefficients $\xi^{(i)}$. The marginal PDF of the identified random variable is obtained by plotting the kernel density estimation of 10000 PC realizations of ξ . These PC realizations are calculated by generating 10000 standard normal random vectors and inserting them into Eq. (28). It is seen that a $p = 10$ order expansion is enough to accurately identify the random coefficients. It is also interesting to note that the PC expansion converges slowly for the first two random coefficients ξ_1 and ξ_2 . This is because more variance is contained in the larger eigenvectors due to the property of PCA.

Some of the realizations generated through sampling the PC representation of linear PCA are shown in Fig. 10. The failure of linear PCA is more pronounced in this case. The realizations definitely do not reflect the original channelized structure of the permeability field. In addition, the permeability value is not correctly predicted.

Similar results are shown in Figs. 11 and 12 for Kernel PCA. However, using Kernel PCA, it is seen that the generated realizations clearly show channelized geological structure with correct permeability values.

We introduce the relative errors for statistics of the experimental samples, $\{\mathbf{y}_i\}_{i=1}^{1000}$, and PC realizations, $\{\mathbf{y}_i^{(\text{pc})}\}_{i=1}^{n_{\text{pc}}}$:

$$e = \frac{\sqrt{\sum_{i=1}^{2025} (T_i - T_i^{(\text{pc})})^2}}{\sqrt{\sum_{i=1}^{2025} T_i^2}}, \quad (48)$$

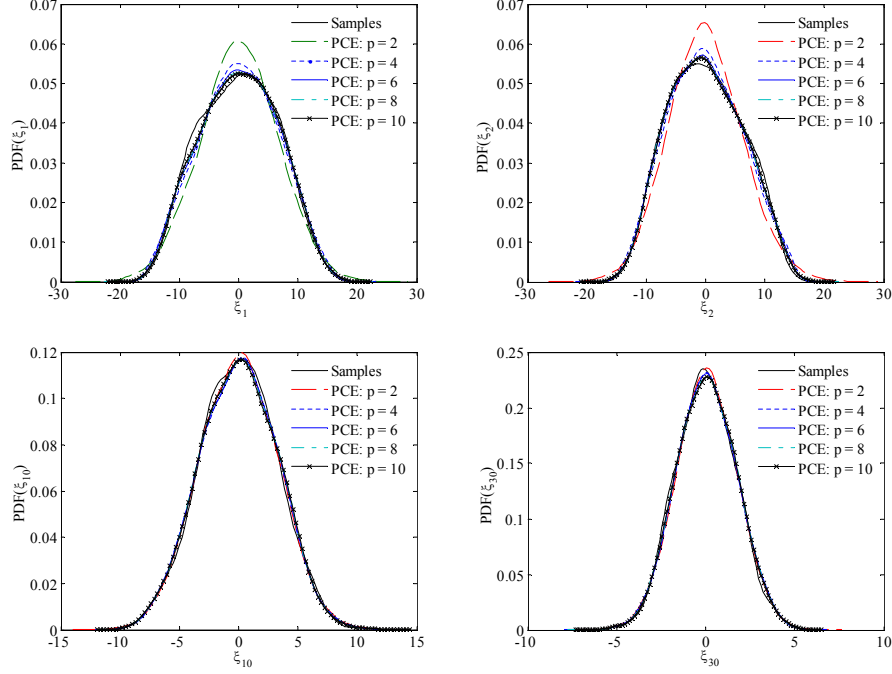


Fig. 9. Four different marginal PDFs of the initial and identified random variables using linear PCA with Hermite chaos.

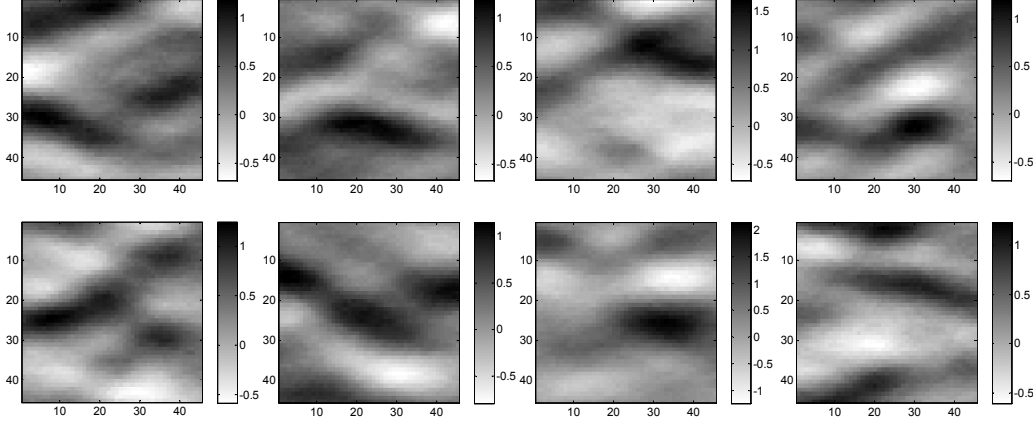


Fig. 10. Realizations of the random field by sampling the corresponding PC representation using linear PCA.

in which T_i represents the appropriate sample statistics of experimental samples $\{\mathbf{y}_i\}_{i=1}^{1000}$ in i th grid block and $S_i^{(\text{pc})}$ represents the corresponding sample statistics of PC realizations $\{\mathbf{y}_i^{(\text{pc})}\}_{i=1}^{n_{\text{pc}}}$. 10000 PC realizations of the random permeability are generated. The results are shown in Table 1. From the table, as expected, the linear PCA gives accurate mean and covariance, which are first- and second-order statistics. On the other hand, the Kernel PCA gives better skewness and kurtosis, which can be considered as third- and fourth-order statistics, respectively. It is noted that relative error of skewness is 1.006 for linear PCA, while it is 0.114 for Kernel PCA. Therefore, the skewness is the dominant higher-order statistic in the case of channelized permeability field.

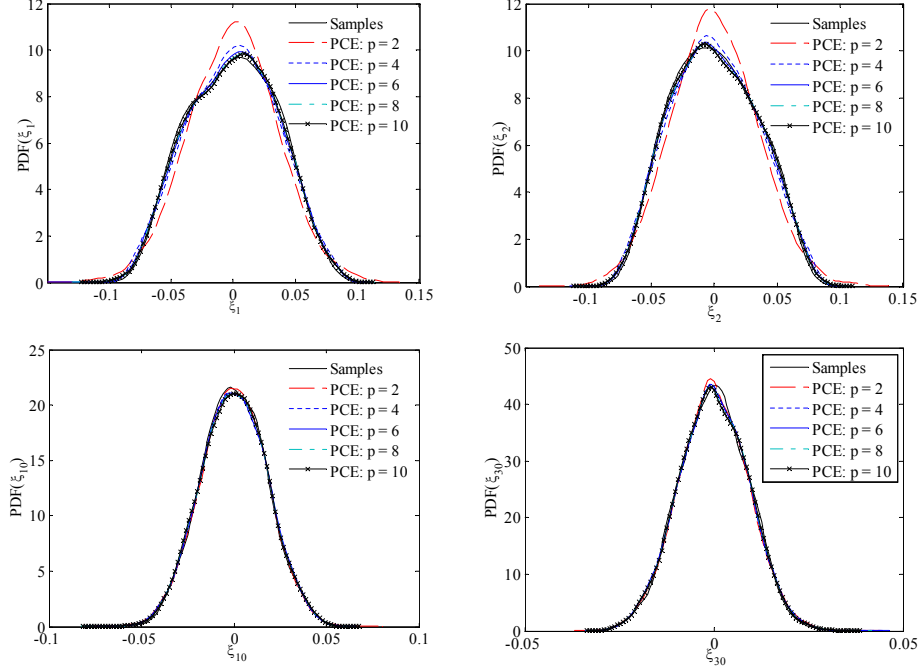


Fig. 11. Four different marginal PDFs of the initial and identified random variables using Kernel PCA with Hermite chaos.

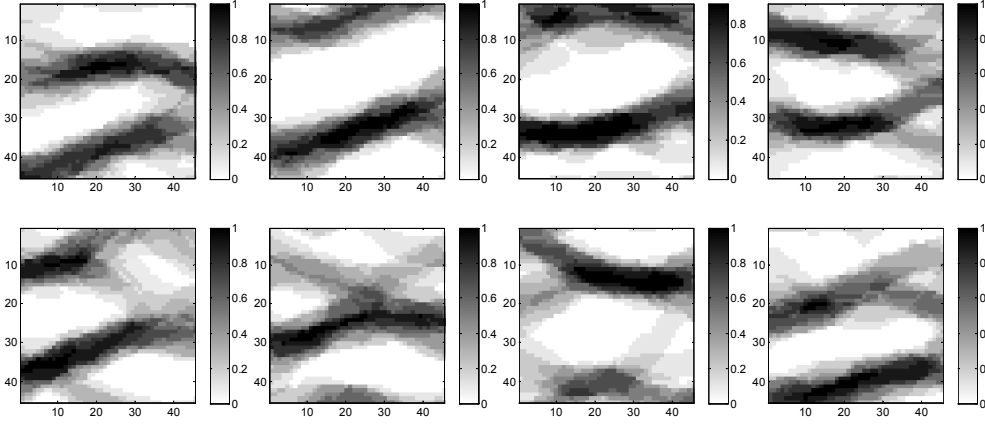


Fig. 12. Realizations of the random field by sampling the corresponding PC representation using Kernel PCA.

Finally, we want to comment on the overall statistics of generated PC realizations. From Table 1, it is seen that the errors from Kernel PCA are still big for covariance and kurtosis. However, it is emphasized that we are most interested in the main geological structure of the random field. We cannot obtain exactly the same statistics as in the experimental samples. There are three possible reasons. First, only an approximate pre-image is calculated. This process contributes significant portion of the total error. Second, a PC expansion is used to fit the marginal PDF of the random coefficients. From Fig. 11, it is noted that some experimental samples are from the long tail region of the marginal PDF. Thus, the experimental samples can be considered as the extreme real-

izations of the underlying random fields. When we sample the corresponding marginal PDF from PC expansion, it is clear that most of the samples are from the high density region. Third, there are only 1000 data samples used to compute the statistics, which may be not enough to get converged statistics. Therefore, the statistics of the generated PC realizations are expected to be slightly smaller than that of experimental sample. But the PC realizations should capture the main statistical features of the experimental samples. This point is demonstrated in the next section.

Table 1

Comparison of statistics between experimental samples and PC samples of the random permeability field.

Errors	Mean vector	Covariance matrix	Skewness vector	Kurtosis vector
Linear PCA	0.039	0.193	1.006	0.676
Kernel PCA	0.102	0.427	0.114	0.496

5.4 Forward uncertainty propagation with the stochastic input model

In this section, the generated stochastic input model is used as an input to the single-phase flow problem on the domain $[0, 1]^2$. The governing equations of the single-phase flow are [43]

$$\nabla \cdot \mathbf{u}(\mathbf{x}, \omega) = 0, \quad \mathbf{u}(\mathbf{x}, \omega) = -a(\mathbf{x}, \omega) \nabla p(\mathbf{x}, \omega) \quad \forall \mathbf{x} \in D, \quad (49)$$

$$\frac{\partial S(\mathbf{x}, t, \omega)}{\partial t} + \mathbf{u}(\mathbf{x}, t, \omega) \cdot \nabla S(\mathbf{x}, t, \omega) = 0, \quad \forall \mathbf{x} \in D, t \in [0, T]. \quad (50)$$

Flow is induced from left-to-right with Dirichlet boundary conditions $p = 1$ on $\{x_1 = 0\}$, $p = 0$ on $\{x_1 = 1\}$ and no-flow homogeneous Neumann boundary conditions on the other two edges. We also impose zero initial condition for saturation $S(x, 0) = 0$ and boundary condition $S(x, t) = 1$ on the in-flow boundary $\{x_1 = 0\}$. Mixed finite element method developed in [43] is used to solve the above equations with spatial discretization 45×45 . So the permeability is defined as a constant in each grid block.

The stochastic permeability is $a = \exp(\mathbf{y})$ if the experimental samples are used or $a = \exp(\mathbf{y}^{n_{pc}})$ if the PC realizations from the stochastic input model are used in the computation. Monte Carlo (MC) simulation is then conducted using both 1000 experimental samples directly and generated PC realizations. 50000 PC realizations are generated by sampling the low-dimensional space. The stochastic input model is the same as the one used to generate the realizations in the last section.

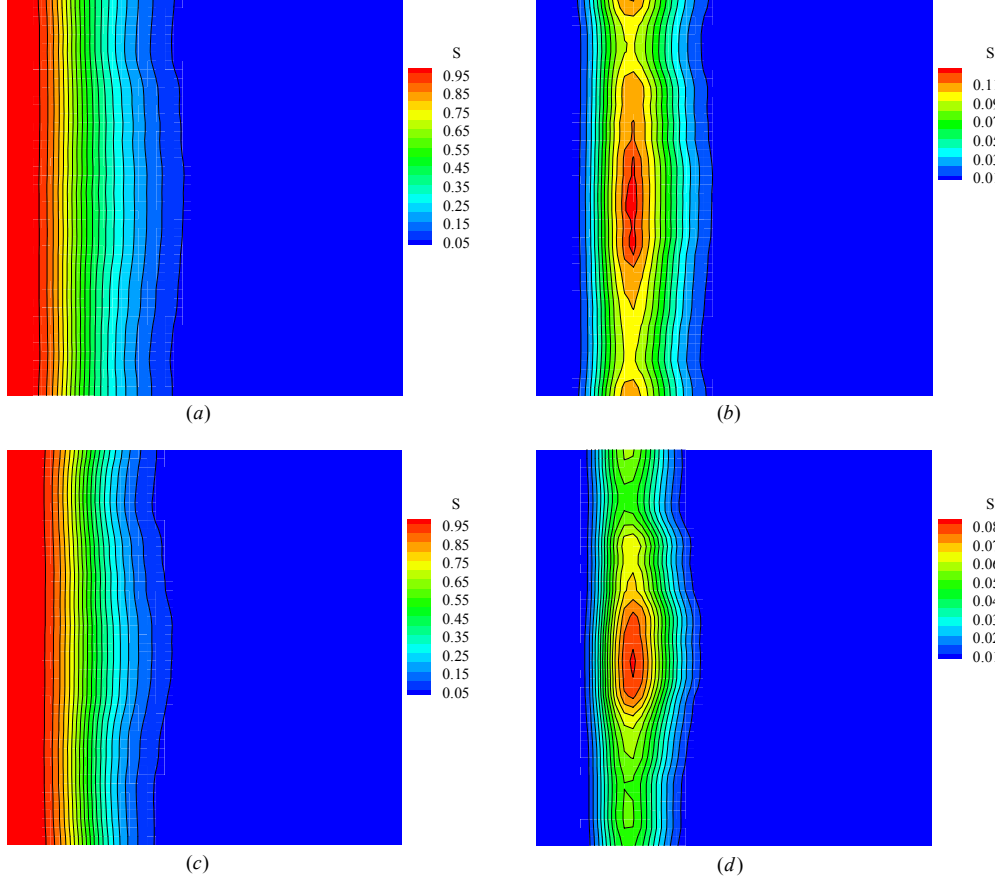


Fig. 13. Contour of saturation at 0.2 PVI: MC mean (a) and variance (b) from experimental samples; MC mean (c) and variance (d) from PC realizations.

The contour plots of saturation at 0.2 PVI are given in Fig. 13. PVI represents dimensionless time and is computed as $PVI = \int Q dt / V_p$, where V_p is the total pore volume of the system, which is equal to the area of the domain D here and $Q = \int_{\partial D^{\text{out}}} (\mathbf{u}_h \cdot \mathbf{n}) ds$ is the total flow rate on the out flow boundary ∂D^{out} . The water cut curves are also given in Fig. 14. The water-cut is defined as $F(t) = \frac{\int_{\partial D^{\text{out}}} (\mathbf{u}_h \cdot \mathbf{n}) S ds}{\int_{\partial D^{\text{out}}} (\mathbf{u}_h \cdot \mathbf{n}) ds}$. From the figures, it is seen that we obtain nearly the same mean saturation and mean water cut curves from the experimental samples and the PC realizations. However, the variances of saturation and water curve curves are not exactly the same. In general, the value of the variance from PC realizations is smaller than that from the experimental samples. But it should be noted that the results from the PC realizations capture the same main features as the results from the experimental samples. To be specific, the regions with the highest variance values of saturation are nearly the same in both cases. The shapes of the variance of the flow curve are also the same in both cases. These observations are expected. As explained before, the experimental samples can be considered as the extreme realizations of the stochastic model. On the other hand, most of the PC realizations are from the highest probability density region of the same stochastic space. Together

with the fact that there are only 1000 experimental samples, it is obvious that the MC results from experimental data are far from converged results. If more experimental samples available, the results are expected to converge to the results using our low-dimensional model. This is possibly expensive in practical engineering problems. On the other hand, our proposed stochastic input model provides a fast way to generate many realizations, which are consistent, in a useful sense, with the experimental data.

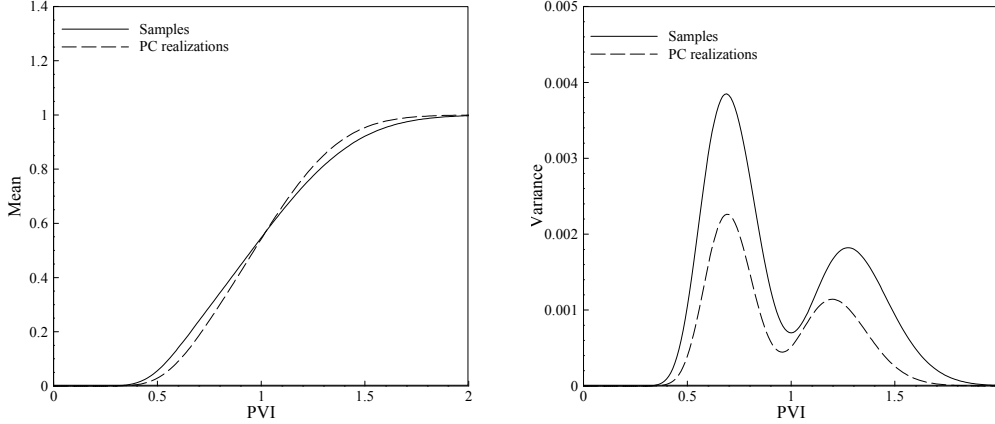


Fig. 14. Comparison of water cut curves: (a) Mean; (b) Variance.

Finally, we want to comment on the computational time. To generate the 1000 experimental data using the ‘snesim’ algorithm, it took approximate 30 minutes on a single processor machine. On the other hand, to generate 50000 PC realizations from our reduced-order model, less than 1 minute is needed on the same machine. In addition, the process of generating PC realizations is easily parallelized, which is used in our MC simulations.

6 Conclusion

In this paper, a new parametric stochastic reduced-order modeling method is proposed, which can efficiently preserve higher-order statistics of the random field given limited experimental data. This method relies on the theory of Kernel PCA to transform the data into a feature space, in which the data is more linearly related than in the original input space. PC expansion is then utilized to identify the random coefficients in the Kernel PCA formulation. A simple but accurate pre-image algorithm is also introduced to project the generated PC realizations back to the original input space. A thorough comparison between the Linear PCA and Kernel PCA on reduced-order modeling of the channelized permeability field is conducted. As expected, Kernel PCA gives more accurate realizations which reflects the original channelized geological structure with much less eigenvectors retained. This results in a low-dimensional stochastic input space. On the other hand, linear PCA fails

to capture the channelized structure with the same number of eigenvectors due to the nonlinearity of the data. Forward uncertainty quantification is also conducted which shows the introduced stochastic input model indeed captures the main statistical properties of the underlying random field.

As a first step towards the implementation of this method, the independence between the random coefficients was assumed. Although this gives us meaningful results in the numerical examples considered, the effect taking the dependence into account as in [19] using Rosenblatt transformation needs further investigation. In addition, the Euclidean distances between data is directly used as the distance measure between data. As in [37], it will be more interesting to take distances between flow responses as the distance measure between permeability data. This is expected to be more accurate since it incorporates the information of the underlying stochastic system.

The basic model reduction ideas envisioned in this work are not limited to generation of viable stochastic input models for property variations. This framework has direct applicability to inverse problems [13], where the generated model can be considered as the prior model of all available properties. Thus, finding the unknown property is only limited to the low-dimensional space and is expected to be more efficient than working in the original high-dimensional space. Furthermore, the generation of a low-dimensional surrogated space provides a prerequisite in stochastic low-dimensional modeling of SPDEs [44], which may have major ramifications in design under uncertainty and stochastic optimization problems. These potentially exciting areas of application of our proposed framework offer fertile avenues of further research.

Acknowledgements

This research was supported by the U.S. Department of Energy, Office of Science, Advanced Scientific Computing Research, the Computational Mathematics program of the National Science Foundation (NSF) (award DMS-0809062) and an OSD/AFOSR MURI09 award on uncertainty quantification. The computing for this research was supported by the NSF through TeraGrid resources provided by NCSA under grant number TG-DMS090007.

References

- [1] D. Zhang, Stochastic method for flow in porous media: coping with uncertainties, Academic Press, San Diego, 2002.

- [2] D. Zhang, Z. Lu, An efficient, high-order perturbation approach for flow in random porous media via karhunen-love and polynomial expansions, *Journal of Computational Physics* 194 (2004) 773 – 794.
- [3] R. Ghanem, P. D. Spanos, *Stochastic Finite Elements: A Spectral Approach*, Springer - Verlag, New York, 1991.
- [4] D. Xiu, G. E. Karniadakis, The Wiener–Askey Polynomial Chaos for stochastic differential equations, *SIAM Journal on Scientific Computing* 24 (2002) 619–644.
- [5] D. Xiu, G. E. Karniadakis, Modeling uncertainty in flow simulations via generalized polynomial chaos, *Journal of Computational Physics* 187 (2003) 137 – 167.
- [6] X. Ma, N. Zabararas, A stabilized stochastic finite element second-order projection method for modeling natural convection in random porous media, *Journal of Computational Physics* 227 (2008) 8448 – 8471.
- [7] I. Babuška, F. Nobile, R. Tempone, A stochastic collocation method for elliptic partial differential equations with random input data, *SIAM Journal on Numerical Analysis* 45 (2007) 1005–1034.
- [8] D. Xiu, J. S. Hesthaven, High-order collocation methods for differential equations with random inputs, *SIAM Journal on Scientific Computing* 27 (2005) 1118–1139.
- [9] B. Ganapathysubramanian, N. Zabararas, Sparse grid collocation schemes for stochastic natural convection problems, *Journal of Computational Physics* 225 (2007) 652 – 685.
- [10] F. Nobile, R. Tempone, C. G. Webster, A sparse grid stochastic collocation method for partial differential equations with random input data, *SIAM Journal on Numerical Analysis* 46 (2008) 2309–2345.
- [11] X. Ma, N. Zabararas, An adaptive hierarchical sparse grid collocation algorithm for the solution of stochastic differential equations, *Journal of Computational Physics* 228 (2009) 3084 – 3113.
- [12] X. Ma, N. Zabararas, An adaptive high-dimensional stochastic model representation technique for the solution of stochastic partial differential equations, *Journal of Computational Physics* 229 (2010) 3884 – 3915.
- [13] X. Ma, N. Zabararas, An efficient bayesian inference approach to inverse problems based on an adaptive sparse grid collocation method, *Inverse Problems* 25 (2009) 035013.
- [14] G. Lin, A. Tartakovsky, An efficient, high-order probabilistic collocation method on sparse grids for three-dimensional flow and solute transport in randomly heterogeneous porous media, *Advances in Water Resources* 32 (2009) 712 – 722.
- [15] M. Loève, *Probability Theory*, fourth ed., Berlin: Springer-Verlag, 1977.

- [16] I. T. Jolliffe, Principal component analysis, second ed., New York: Springer-Verlag, 2002.
- [17] C. Desceliers, R. Ghanem, C. Soize, Maximum likelihood estimation of stochastic chaos representations from experimental data, *International Journal for Numerical Methods in Engineering* 66 (2006) 978 – 1001.
- [18] R. G. Ghanem, A. Doostan, On the construction and analysis of stochastic models: Characterization and propagation of the errors associated with limited data, *Journal of Computational Physics* 217 (2006) 63 – 81.
- [19] S. Das, R. Ghanem, J. C. Spall, Asymptotic sampling distribution for polynomial chaos representation from data: A maximum entropy and fisher information approach, *SIAM Journal on Scientific Computing* 30 (2008) 2207–2234.
- [20] S. Das, R. Ghanem, S. Finette, Polynomial chaos representation of spatio-temporal random fields from experimental measurements, *Journal of Computational Physics* 228 (2009) 8726 – 8751.
- [21] G. Stefanou, A. Nouy, A. Clement, Identification of random shapes from images through polynomial chaos expansion of random level set functions, *International Journal for Numerical Methods in Engineering* 79 (2009) 127 – 155.
- [22] M. Arnst, R. Ghanem, C. Soize, Identification of Bayesian posteriors for coefficients of chaos expansions, *Journal of Computational Physics* 229 (2010) 3134 – 3154.
- [23] C. Soize, Identification of high-dimension polynomial chaos expansions with random coefficients for non-gaussian tensor-valued random fields using partial and limited experimental data, *Computer Methods in Applied Mechanics and Engineering* 199 (2010) 2150 – 2164.
- [24] M. Rosenblatt, Remarks on a multivariate transformation, *Ann. Math. Statist.* 23 (1952) 470 – 472.
- [25] K. Sargsyan, B. Debusschere, H. Najm, O. L. Maître, Spectral representation and reduced order modeling of the dynamics of stochastic reaction networks via adaptive data partitioning, *SIAM Journal on Scientific Computing* 31 (2010) 4395–4421.
- [26] I. Babuška, K.-M. Liu, R. Tempone, Solving stochastic partial differential equations based on the experimental data, *Mathematical Models and Methods in Applied Sciences* 13 (2003) 415 – 444.
- [27] B. Ganapathysubramanian, N. Zabaras, Modeling diffusion in random heterogeneous media: Data-driven models, stochastic collocation and the variational multiscale method, *Journal of Computational Physics* 226 (2007) 326 – 353.
- [28] N. Agarwal, N. R. Aluru, A data-driven stochastic collocation approach for uncertainty quantification in mems, *International Journal for Numerical Methods in Engineering* 83 (2010) 575 – 597.

- [29] B. Ganapathysubramanian, N. Zabaras, A non-linear dimension reduction methodology for generating data-driven stochastic input models, *Journal of Computational Physics* 227 (2008) 6612 – 6637.
- [30] J. B. Tenenbaum, V. d. Silva, J. C. Langford, A Global Geometric Framework for Nonlinear Dimensionality Reduction, *Science* 290 (2000) 2319–2323.
- [31] B. Scholkopf, A. Smola, K.-R. Muller, Nonlinear component analysis as a kernel eigenvalue problem, *Neural Computation* 10 (1998) 1299–1319.
- [32] S. Bernhard, S. Alexander, *Learning With Kernels*, MIT Press, 2002.
- [33] J. Shawe-Taylor, N. Cristianini, *Kernel Methods for Pattern Analysis*, Cambridge University Press, 2004.
- [34] S. Mika, S. Bernhard, S. Alexander, M. Klaus-Robert, M. Scholz, G. Ratsch, Kernel PCA and de-noising in feature spaces, in: *Advances in Neural Information Processing Systems 11*, MIT Press, 1999, pp. 536–542.
- [35] Y. Rathi, S. Dambreville, A. Tannenbaum, Statistical shape analysis using kernel PCA, in: *Image Processing: Algorithms and Systems, Neural Networks, and Machine Learning*, SPIE, 2006, p. 60641B.
- [36] P. Sarma, L. J. Durlofsky, K. Aziz, Kernel principal component analysis for efficient, differentiable parameterization of multipoint geostatistics, *Mathematical Geosciences* 40 (2008) 3–32.
- [37] C. Scheidt, J. Caers, Representing spatial uncertainty using distances and kernels, *Mathematical Geosciences* 41 (2009) 397–419.
- [38] J.-Y. Kwok, I.-H. Tsang, The pre-image problem in kernel methods, *Neural Networks, IEEE Transactions on* 15 (2004) 1517–1525.
- [39] S. S. Ravindran, A reduced-order approach for optimal control of fluids using proper orthogonal decomposition, *International Journal for Numerical Methods in Fluids* 34 (2000) 425–448.
- [40] A. W. Bowman, A. Azzalini, *Applied Smoothing Techniques for Data Analysis: the kernel approach with S-Plus illustrations*, Oxford University Press, 1997.
- [41] C. K. Williams, On a connection between kernel PCA and metric multidimensional scaling, *Machine Learning* 46 (2002) 11–19.
- [42] S. Strebelle, Conditional simulation of complex geological structures using multiple-point statistics, *Mathematical Geology* 34 (2002) 1–21.
- [43] X. Ma, N. Zabaras, A stochastic mixed finite element heterogeneous multiscale method for flow in porous media, *Journal of Computational Physics*, Submitted.
- [44] D. Venturi, X. Wan, G. E. Karniadakis, Stochastic low-dimensional modelling of a random laminar wake past a circular cylinder, *Journal of Fluid Mechanics* 606 (2008) 339–367.