



Structural reliability analysis based on ensemble learning of surrogate models

Kai Cheng, Zhenzhou Lu*

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

ARTICLE INFO

Keywords:

Ensemble learning
Surrogate model
Reliability analysis
Active learning

ABSTRACT

Assessing the failure probability of complex structure is a difficult task in presence of various uncertainties. In this paper, a new adaptive approach is developed for reliability analysis by ensemble learning of multiple competitive surrogate models, including Kriging, polynomial chaos expansion and support vector regression. Ensemble of surrogates provides a more robust approximation of true performance function through a weighted average strategy, and it helps to identify regions with possible high prediction error. Starting from an initial experimental design, the ensemble model is iteratively updated by adding new sample points to regions with large prediction error as well as near the limit state through an active learning algorithm. The proposed method is validated with several benchmark examples, and the results show that the ensemble of multiple surrogate models is very efficient for estimating failure probability ($> 10^{-4}$) of complex system with less computational costs than the traditional single surrogate model.

1. Introduction

Structural reliability analysis is of great importance in engineering, it aims at computing the probability of failure of a system with respect to some performance criterion in the presence of various uncertainties. For a given structural system with n -dimensional input parameter $\mathbf{x} = (x_1, x_2, \dots, x_n)^T$, the performance function (also known as limit state function) $g(\mathbf{x})$ divides the input variable space into two domains, i.e. the safety domain ($g(\mathbf{x}) > 0$) and failure domain ($g(\mathbf{x}) \leq 0$). Thus the failure probability P_f reads:

$$P_f = \int_{I_{g(\mathbf{x}) \leq 0}(\mathbf{x})} f_{\mathbf{x}}(\mathbf{x}) d\mathbf{x}, \quad (1)$$

where $I_{g(\mathbf{x}) \leq 0}(\mathbf{x}) = \begin{cases} 1 & g(\mathbf{x}) \leq 0 \\ 0 & g(\mathbf{x}) > 0 \end{cases}$ is the indicator function of the failure domain and $f_{\mathbf{x}}(\mathbf{x})$ is the joint probability density function (PDF) of $\mathbf{x}$.

The task of reliability analysis is to perform the integration in Eq. (1). Generally, the reliability analysis method in literature can be categorized into three types: approximate analytical methods [1,2], numerical integration methods [3–7] and numerical simulation methods [8–16]. The approximate analytical methods expand the performance function $g(\mathbf{x})$ at mean point or design point by Taylor expansion, and ignore the higher order terms to estimate the failure probability, such as First order reliability method (FORM) [1] and Second order reliability analysis method (SORM) [2]. However, their accuracy can hardly be

guaranteed especially for highly non-linear problems. For numerical integration methods, the first few moments of the performance function are computed by the point estimation method [3–5] or sparse grid integration method [6], and then failure probability is estimated by these moments. However, the computation cost of these methods increases sharply (exponentially) with the input variable dimensionality. The numerical simulation methods include Monte Carlo simulation (MCS), Importance sampling (IS) [8,9,17], Subset simulation (SS) [11–13,18] and recent work called thermodynamic integration and parallel tempering (TIPT) method [19]. These methods are relatively robust to the type and dimension of the problem, but they cannot satisfy the computational efficiency requirements for time-consuming model. To reduce the computational cost for reliability analysis, surrogate-assisted methods have received much attention in the past few decades. These methods aim at constructing a surrogate model (also known as *meta-model*) with an explicit expression based on a set of observed points to approximate the true performance function, and thus one can perform reliability analysis efficiently based on the cheap-to-evaluate surrogate model. For decades, several types of surrogate models are available in literatures for reliability analysis including polynomial chaos expansion (PCE) [20], Kriging (Gaussian Process, GP) [21–24], support vector machine (SVM) and support vector regression (SVR) [12,25,26], high dimensional model representation (HDMR) [27] and so on.

Recently, surrogate models combined with active learning strategies

* Corresponding author.

E-mail address: zhenzhoulu@nwpu.edu.cn (Z. Lu).

have been well developed for reliability analysis [21,23,24,28–34] of complex system, especially for Kriging model. These methods start from an initial design of experiment (DoE), and enrich it sequentially by adding new sample points based on the predefined learning function. The learning function is usually developed based on the statistical information of surrogate model from different perspectives. The expected feasibility function (EFF) proposed by Bichon et al. [35] and U function developed by Echard et al. [28,29] both select points near the limit state surface of Kriging model with large prediction uncertainty. The expected risk function (ERF) in Ref [36] and H function in Ref [21] search for points with large prediction error and information entropy in the vicinity of limit state surface of Kriging model respectively. Sun et al. [23] developed the least improvement function (LIF) based on Kriging, which quantifies the improvement of the accuracy of estimated failure probability when new points are added to the DoE. Marelli et al. [20] presented a learning function that focuses on the probability of misclassification of PCE model based on bootstrap resampling strategy. In Ref [37], adaptive SVR model is presented to estimate the failure probability, where the learning function is defined by the distance criteria. It has been proven that these well-developed learning functions are very efficient to improve the accuracy and efficiency for reliability analysis of complex models.

In this paper, instead of fitting the performance function with single surrogate model, we explore the possibility of ensemble of multiple competitive surrogate models to approximate the performance function for reliability analysis. Each surrogate model is required to predict the model response, and the final prediction is obtained by a weighted average of multiple surrogate models. In the meanwhile, the local prediction error is estimated by the variance of the multiple surrogate models, thus an active learning algorithm is developed to select some dangerous sample points sequentially in the regions with large prediction error to improve the prediction accuracy as much as possible.

The layout of this paper is as follows. Section 2 presents an overview of PCE, Kriging and SVR surrogate models. In section 3, the proposed active learning strategy for reliability analysis is presented. Four examples are employed in Section 4 to demonstrate the efficiency and accuracy of the proposed method. Finally, some conclusions are drawn in Section 5.

2. Review of surrogate models

In this section, three kinds of surrogate models, namely, PCE, Kriging and SVR are briefly reviewed.

2.1. Polynomial chaos expansion

The classic PCE was first proposed by Wiener [38,39] in the 30s. The key concept of PCE is to expand the model response onto basis made of multivariate polynomials that are orthogonal with respect to the joint distribution of the input variables. In this setting, characterizing the response probability density function (PDF) is equivalent to evaluate the PC coefficients, i.e. the coordinates of the random response in this basis. The classic PCE of order p for n -dimensional random variable $\mathbf{x}$ can be expressed as [40]:

$$\tilde{g}_p(\mathbf{x}) = \sum_{0 \leq |\alpha| \leq p} \omega_\alpha \psi_\alpha(\mathbf{x}), \quad |\alpha| = \sum_{i=1}^n \alpha_i, \quad (2)$$

where $\alpha = \{\alpha_1, \dots, \alpha_n\} (\alpha_i \geq 0)$ is the multidimensional index notation vector, ω_α is the unknown deterministic coefficients vector, and $\psi_\alpha(\mathbf{x})$ is the multivariate polynomial vector. The total number of the expansion terms in the summation of Eq. (2) is $P = (p + n)!/p!n! + 1$. To calculate the PCE coefficients, the traditional projection method and regression method [41,42] suffer from the so-called *curse of dimensionality*.

To overcome this issue, many attempts have been made to develop sparse PCE model in the field of uncertainty quantification (UQ) [43–50], the common idea holding in these methods is that the PCE

coefficients are sparse (i.e. having only several dominant coefficients). Given the training sample $\{\mathbf{X}, \mathbf{Y}\}$, where $\mathbf{X} = \{\mathbf{x}_1, \dots, \mathbf{x}_N\}^T$ are the input data, $\mathbf{Y} = \{Y_1, \dots, Y_N\}^T$ are the corresponding model responses and N is the size of sample, the dominant PCE coefficients can be recovered by solving the following optimization problem

$$\omega_\alpha = \arg \min_{\omega_\alpha} \|\omega_\alpha\|_1 \quad \text{subject to} \quad \|\Phi \omega_\alpha - \mathbf{Y}\| \leq \delta \quad (3)$$

where $\|\omega_\alpha\|_1$ is the l_1 norm of PCE coefficients, δ is a tolerance parameter of the truncation error and $\Phi = [\psi_j(\mathbf{x}_i)]_{ij} \in \mathbb{R}^{N \times P}$ is the measure matrix. In this paper, the least angle regression (LAR) technique is used to develop sparse PCE, which is available from the UQLab toolbox [51].

2.2. Kriging/Gaussian process

Kriging model (also known as GP) was first proposed in the field of geostatistics by Krige [52] and Matheron [53]. It tends to find the best linear unbiased predictor while minimizes the mean square error of the prediction. The universal Kriging is composed of a polynomial term used for global trend prediction and a Gaussian process term used for local deviation regression, which can be expressed as

$$g_K(\mathbf{x}) = p^T(\mathbf{x})\boldsymbol{\beta} + Z(\mathbf{x}), \quad (4)$$

where $p(\mathbf{x}) = [p_1(\mathbf{x}), \dots, p_M(\mathbf{x})]^T$ is the polynomial basis function, $\boldsymbol{\beta} = [\beta_1, \dots, \beta_M]^T$ represents the corresponding regression coefficient vector, and $Z(\mathbf{x})$ is a Gaussian process with zero mean and covariance defined as

$$\text{Cov}(Z(\mathbf{x}_i), Z(\mathbf{x}_j)) = \sigma^2 R(\mathbf{x}_i, \mathbf{x}_j, \boldsymbol{\theta}), \quad (5)$$

where σ^2 is the variance of $Z(\mathbf{x})$, and $R(\mathbf{x}_i, \mathbf{x}_j, \boldsymbol{\theta})$ is the correlation coefficient between $Z(\mathbf{x}_i)$ and $Z(\mathbf{x}_j)$ with parameters $\boldsymbol{\theta} = [\theta_1, \dots, \theta_n]^T$. The correlation function controls the smoothness of the Kriging model, and here the Gaussian correlation function is used in the present work, which is defined as

$$R(\mathbf{x}_i, \mathbf{x}_j, \boldsymbol{\theta}) = \prod_{k=1}^n \exp[-\theta_k (x_i^{(k)} - x_j^{(k)})^2]. \quad (6)$$

Generally, the unknown hyper-parameters $\boldsymbol{\gamma} = (\boldsymbol{\beta}, \sigma^2, \boldsymbol{\theta})$ of Kriging model can be tuned by maximum likelihood estimation technique. After the optimal parameters $\hat{\boldsymbol{\gamma}} = (\hat{\boldsymbol{\beta}}, \hat{\sigma}^2, \hat{\boldsymbol{\theta}})$ are obtained, the posterior distribution of $g_K(\mathbf{x})$ is a Gaussian process $g_K(\mathbf{x}) \sim \mathcal{N}(\mu(\mathbf{x}), \sigma^2(\mathbf{x}))$ with mean

$$\tilde{g}_K(\mathbf{x}) = \mu(\mathbf{x}) = \mathbf{p}^T(\mathbf{x})\hat{\boldsymbol{\beta}} + \mathbf{r}^T(\mathbf{x})\mathbf{R}^{-1}(\mathbf{Y} - \mathbf{F}\hat{\boldsymbol{\beta}}), \quad (7)$$

and variance

$$\begin{aligned} \sigma^2(\mathbf{x}) &= \hat{\sigma}^2 [1 - \mathbf{r}^T(\mathbf{x})\mathbf{R}^{-1}\mathbf{r}(\mathbf{x}) + (\mathbf{r}(\mathbf{x})\mathbf{R}^{-1}\mathbf{F} - \mathbf{p}(\mathbf{x}))^T (\mathbf{F}^T\mathbf{R}^{-1}\mathbf{F})^{-1} \\ &\quad (\mathbf{r}(\mathbf{x})\mathbf{R}^{-1}\mathbf{F} - \mathbf{p}(\mathbf{x}))], \end{aligned} \quad (8)$$

where $\mathbf{r}(\mathbf{x}) = [R(\mathbf{x}, \mathbf{x}_1), \dots, R(\mathbf{x}, \mathbf{x}_N)]^T$ represents the correlation vector between $\mathbf{x}$ and N observed points, $\mathbf{R} = [R(\mathbf{x}_i, \mathbf{x}_j)]_{ij} \in \mathbb{R}^{N \times N}$ is the correlation matrix and $\mathbf{F} = [\mathbf{f}_j(\mathbf{x}_i)]_{ij} \in \mathbb{R}^{N \times M}$. The posterior mean in Eq. (7) is known as the Kriging predictor $\tilde{g}_K(\mathbf{x})$. The posterior variance formula of Eq. (8) corresponds to the mean squared error (MSE) of this predictor and it is also known as the Kriging variance.

2.3. Support vector regression

Support-vector regression (SVR) was developed on statistical learning theory by Vapnik [54,55]. Generally, a linear SVR model is formulated as:

$$\tilde{g}_S(\mathbf{x}) = \omega \cdot \mathbf{x} + b, \quad (9)$$

where $\omega \in \mathbb{R}^n$ is the coefficient vector and $b \in \mathbb{R}$ is a constant. The goal of SVR is to find a function $\tilde{g}_S(\mathbf{x})$ that can estimate the output response value whose deviation is less than ε from the real targets of the training

Download English Version:

<https://daneshyari.com/en/article/13422611>

Download Persian Version:

<https://daneshyari.com/article/13422611>

[Daneshyari.com](https://daneshyari.com)