



An improved adaptive kriging-based importance technique for sampling multiple failure regions of low probability

F. Cadini^{a,*}, F. Santos^{a,b}, E. Zio^{a,c}

^a Dipartimento di Energia, Politecnico di Milano, Italy

^b Departamento de Energía Nuclear, Politécnica de Madrid, Spain

^c Chair on Systems Science and Energetic Challenge, European Foundation for New Energy-Electricité de France, Ecole Centrale Paris and Supélec, Paris, France

ARTICLE INFO

Article history:

Received 7 March 2014

Received in revised form

6 June 2014

Accepted 27 June 2014

Available online 4 July 2014

Keywords:

Small failure probabilities

Multiple failure regions

Adaptive kriging

Importance sampling

ABSTRACT

The estimation of system failure probabilities may be a difficult task when the values involved are very small, so that sampling-based Monte Carlo methods may become computationally impractical, especially if the computer codes used to model the system response require large computational efforts, both in terms of time and memory. This paper proposes a modification of an algorithm proposed in literature for the efficient estimation of small failure probabilities, which combines FORM to an adaptive kriging-based importance sampling strategy (AK-IS). The modification allows overcoming an important limitation of the original AK-IS in that it provides the algorithm with the flexibility to deal with multiple failure regions characterized by complex, non-linear limit states. The modified algorithm is shown to offer satisfactory results with reference to four case studies of literature, outperforming in general several other alternative methods of literature.

© 2014 Elsevier Ltd. All rights reserved.

1. Introduction

Given a probabilistic model, described by a n -dimensional random vector $\mathbf{x} = \{x_1, \dots, x_n\}$ with probability density function (pdf) $f(\mathbf{x})$, and a performance function $G(\mathbf{x})$ representing the system response, failure is usually defined as the event

$$F = \{G(\mathbf{x}) \leq 0\} \quad (1)$$

where the set of values $\mathbf{x} : G(\mathbf{x}) \leq 0$ and $\mathbf{x} : G(\mathbf{x}) = 0$ are called failure domain and limit state function, respectively.

By introducing the following failure indicator function $I_F(\mathbf{x})$

$$I_F(\mathbf{x}) = I_{\{G \leq 0\}}(\mathbf{x}) = \begin{cases} 1 & G(\mathbf{x}) \leq 0 \\ 0 & G(\mathbf{x}) > 0 \end{cases} \quad (2)$$

the failure probability can, then, be written as

$$p_f = P\{G(\mathbf{x}) \leq 0\} = E[I_F(\mathbf{x})] = \int_{R^n} I_F(\mathbf{x}) f(\mathbf{x}) d\mathbf{x} \quad (3)$$

For example, in structural reliability problems, the performance function $G(\mathbf{x})$ represents the difference between the loads acting on a structure and its resistance [1,2], or, in the performance assessment of a radioactive waste repository, the difference

between a regulatory threshold and the expected radiological dose [3].

A typical approach to the estimation of the failure probability (3) is that of resorting to a standard, or crude, Monte Carlo (MC) scheme, which amounts to (i) sampling N values of the model input random vector $\mathbf{x}$ from $f(\mathbf{x})$ and (ii) running the model in correspondence of each of the N realizations of $\mathbf{x}$ in order to compute the performance function $G(\mathbf{x})$. The failure probability can, then, be estimated by dividing the number of realizations for which $G(\mathbf{x}) \leq 0$ by N .

However, as engineered systems are very reliable, so that their failure is a rare event, the estimation of such probabilities would require a large number of simulations and, possibly, prohibitive computational times by standard MC schemes. This problem is particularly relevant for cases in which the computer codes used to evaluate the performance function $G(\mathbf{x})$ are computationally very intensive, such as those based on complex finite elements models [1,2,4], for example.

Various methods have been proposed in the literature to address this problem: the interested reader may refer to [5,6] for thorough reviews and comparisons of many existing methods. Here, we will briefly recall the general ideas behind the most widely used methods.

The first family of methods commonly used in structural reliability analysis and known as FORM or SORM (first or second order reliability methods), stems from an approximation of the

* Corresponding author. Tel.: +39 02 2399 6355.

E-mail address: francesco.cadini@polimi.it (F. Cadini).

limit state function around the so called “most probable failure point (MPFP)” or “design point”, based on a Taylor series expansion [7,8]. The MPFP is defined as the point on the limit state function which is closer to the mean of the random input vector $\mathbf{x}$. These methods require the computations of the gradient and the Hessian of the limit state function, which are usually performed by a finite difference scheme. The numerical approximation of the limit state functions allows in general fast estimates of the failure probabilities, requiring a very limited number of performance function evaluations by the original model [12]. However, this effectiveness is obtained at the expense of a few important limitations (i) the methods do not allow any quantification of the approximation errors; (ii) in case of complicated, highly non-linear limit state functions, the linear approximation provided by the FORM introduces large estimation errors, only partially reduced if the SORM is used; (iii) in presence of multiple, non-connected failure domains the methods may lead to biased estimates of the failure probability, although some efforts to address this issue have been made in the past [9,10]; and (iv) when dealing with high dimensional input spaces, the finite difference scheme may severely affect the efficiency of these methods.

The second family of methods, also known as simulation methods, comprises those based on MC schemes. Among these, the crude MC approach described above is the simplest and probably the most widely used method for estimating failure probabilities, but, as mentioned before, it becomes very inefficient when dealing with rare events. For this reason, many so called variance reduction techniques have been proposed in literature, which aim at developing more efficient MC estimators achieving the same levels of accuracy at largely reduced numbers of model simulations, i.e. evaluations of the performance function. Perhaps, the most popular variance reduction technique is that of Importance Sampling (IS), which has been successfully applied in many fields of research, e.g. probabilistic risk assessment of industrial systems [3,11], structural reliability [2,12,13], queueing models of telecommunication systems [14,15], project management [16], network reliability [17], etc. In IS, a suitable importance density alternative to the original input pdf $f(\mathbf{x})$ is chosen so as to favor the MC samples to be near the failure region, thus forcing the rare failure event to occur more often [18]. In this regard, it is possible to show that there exists an optimal importance density so that the variance of the MC estimator is zero [18]. Unfortunately, this pdf is not implementable in practice, since its analytical expression depends on the unknown failure probability p_f itself; however, several techniques in various fields of research have been proposed in literature to render the instrumental importance pdf as similar as possible to the optimal one, for obtaining good sampling efficiency. In this respect, a possible approach is, for example, that of the Cross Entropy (CE) method, based on the minimization of the Kullback–Leibler distance between an instrumental pdf belonging to the natural exponential family (NEF) and the optimal one [19]; this method is potentially very attractive, although its limited flexibility prevents its applicability to a wide range of engineering problems. A common approach in structural reliability is that of choosing the importance density as a joint Gaussian distribution centered around the MPFP identified by a FORM (or SORM) in the isoprobabilistically transformed standard input space [12,13]: by doing so, it is possible to refine the result of the FORM (SORM) by an IS procedure which picks the samples in the vicinity of the failure region.

In general, sampling-based methods stemming from a variance reduction technique allows significant improvements with respect to a crude MC simulation; however, they suffer from the fact that the number of time-demanding evaluations of the original performance function required for estimating small probabilities remains too large [2].

The third family of methods for efficiently addressing this problem relies on the substitution of the original performance function by a surrogate model (or metamodel) within a sampling-based scheme; a metamodel is, in general, orders of magnitude faster to be evaluated, thus allowing significant computational savings. Several metamodels have been proposed in literature, such as quadratic response surfaces [20–22], polynomial chaos expansions [23], support vector machines [1,24–26], neural networks [27,28] and kriging [29,30]. The major drawback of the direct substitution of the original performance function with a surrogate model is that it is often impossible to keep the approximation error under control [2].

Recently, adaptive strategies for coupling sampling-based method and metamodeling have been proposed, which allow refining the metamodel construction until a predefined level of accuracy is achieved. For example, [2] proposed to resort to a kriging-based surrogate model to approximate the optimal importance density, thus obtaining a new estimator of the failure probability as the product of a term given by a standard MC estimation based on the kriging approximation, and a correction factor, whose computation is based on a comparison of the outcomes of the original model and the kriging-based metamodel. Following a different strategy, [4] proposed to use kriging to classify a population of candidate points sampled by standard MC from the input uncertainty pdfs, whereby the metamodel training set was, then, iteratively enriched on the basis of a learning function accounting for the probability of the metamodel correct classification (AK-MCS algorithm). The method was also generalized in [31] for estimating failure probabilities in case of systems made up of components connected in some functional logic (parallel, series), each characterized by its own performance function. In order to reduce the computational efforts required by the learning method in case of small failure probability estimation, [12] further improved the AK-MCS method by resorting to IS for sampling the candidate points from an importance density centered about the MPFP previously identified by FORM. This method, called Adaptive Kriging-Importance Sampling (AK-IS), was shown to be very efficient, but the use of FORM limits its application to those problem characterized by a single failure region with a unique MPFP.

In this paper, we propose a modification of the approach proposed by [12], which provides the algorithm with the capability of dealing with multiple, disconnected failure regions. The improvement is based on the replacement of the FORM stage of the original algorithm with the metamodel refinement step of the meta-IS algorithm proposed by Dubourg et al., [2], which is proven to be effective to identify points of failure in the input space, regardless the different failure domains they belong to. A clustering procedure based on the kriging prediction is, then, devised to automatically assign the correct failure region to each of the points previously identified, and select, for each such region, the points closest to the origin, thus including also the approximations of the existing multiple MPFPs. Finally, according to [13,32–34], a multimodal importance joint pdf is centered around these points in the standard space and the efficiency of the second stage of the [12] algorithm can be exploited with a few minor modifications.

The performances of the resulting metaAK-IS² algorithm are verified with reference to four analytic case studies often used in literature [1,4,12,13].

The paper is organized as follows. Section 2 briefly recalls the AK-IS algorithm introduced in [12] (Section 2.1) and the first step of the meta-IS algorithm introduced in [2] (Section 2.2), and presents how the metaAK-IS² scheme can be obtained by combining the two algorithms by means of a K-means-based clustering procedure (Section 2.3). Section 3 illustrates the application of the new algorithm with reference to four analytic case

Download English Version:

<https://daneshyari.com/en/article/806769>

Download Persian Version:

<https://daneshyari.com/article/806769>

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