

# Investigation and benchmark of algorithms for reliability analysis

Jörg F. Unger<sup>1,2</sup> & Dirk Roos<sup>1\*</sup>

<sup>1</sup> DYNARDO – Dynamic Software and Engineering GmbH, Weimar

<sup>2</sup> Institut für Strukturmechanik, Bauhaus-Universität Weimar

## Abstract

The properties of a structural member are in general not static but stochastic parameters due to e.g. irregularities within the production process or the material itself. Furthermore the loads under real conditions are stochastic variables, e.g. wind or the load due to people crossing a bridge. In order to evaluate the failure probability of a structure, the stochastic character of the problem has to be taken into account. Different methods exist, but they all have a limited area of application. The aim of this paper is to investigate the applicability of certain methods for specific problems and to give a general indication, which method is appropriate for certain problem classes.

**Keywords:** reliability, Monte Carlo simulation, latin hypercube, directional sampling, FORM, response surface

---

\*Kontakt: Dr.-Ing. Dirk Roos, DYNARDO – Dynamic Software and Engineering GmbH, Luthergasse 1d, D-99423 Weimar, E-Mail: [dirk.roos@dynardo.de](mailto:dirk.roos@dynardo.de)

# Contents

|          |                                                                        |           |
|----------|------------------------------------------------------------------------|-----------|
| <b>1</b> | <b>Applied methods</b>                                                 | <b>3</b>  |
| 1.1      | General remarks . . . . .                                              | 3         |
| 1.2      | Monte Carlo simulation . . . . .                                       | 3         |
| 1.3      | Latin Hypercube Sampling . . . . .                                     | 4         |
| 1.4      | Adaptive Sampling . . . . .                                            | 4         |
| 1.5      | FORM - First Order Reliability Method . . . . .                        | 6         |
| 1.6      | Importance Sampling Procedure Using Design points . . . . .            | 7         |
| 1.7      | Directional Sampling . . . . .                                         | 8         |
| 1.8      | Adaptive Directional Sampling . . . . .                                | 10        |
| 1.9      | Response surface methods . . . . .                                     | 11        |
| 1.9.1    | Approximation of the limit state function by tangential hyperplanes    | 11        |
| 1.9.2    | Approximation of the limit state function using weighted radii . . .   | 12        |
| 1.9.3    | Calculating the points of support for the limit state function . . . . | 14        |
| <b>2</b> | <b>Limit state functions</b>                                           | <b>14</b> |
| 2.1      | Limit state function 1 . . . . .                                       | 14        |
| 2.2      | Limit state function 2 . . . . .                                       | 17        |
| 2.3      | Limit state function 6 . . . . .                                       | 21        |
| 2.4      | Limit state function 10 . . . . .                                      | 21        |
| 2.5      | Limit state function 11 . . . . .                                      | 22        |
| <b>3</b> | <b>Summary and Conclusion</b>                                          | <b>24</b> |

# 1 Applied methods

## 1.1 General remarks

Given a vector  $\mathbf{X}$  of  $n$  random variables and a limit state function  $g(\mathbf{x})$ , which divides the total domain  $\Omega$  into the failure domain ( $g(\mathbf{x}) \leq 0$ ) and the safe domain ( $g(\mathbf{x}) > 0$ ), the failure probability  $p_f$  is defined as

$$p_f = \int_{g(\mathbf{x}) \leq 0} f(\mathbf{x}) d\mathbf{x}, \quad (1)$$

where  $f(\mathbf{x})$  is the joined probability density function of the random variables  $\mathbf{X}$ . By introducing an indicator function  $I(\mathbf{x})$  with

$$I(\mathbf{x}) = \begin{cases} 1 & \text{if } g(\mathbf{x}) \leq 0 \\ 0 & \text{else} \end{cases} \quad (2)$$

eq.(1) can be reformulated as

$$p_f = \int_{\Omega} I(\mathbf{x}) f(\mathbf{x}) d\mathbf{x}, \quad (3)$$

The analytical evaluation of this integral is in most practical problems impossible. This is due to the fact, that an exact determination of the primitive for many functions  $I(\mathbf{x})f(\mathbf{x})$  is not possible, or that no analytical expression for the limit state function  $g(\mathbf{x})$  can be derived, e.g. if  $g(\mathbf{x})$  is approximated by a limited number of complex FE-solutions.

## 1.2 Monte Carlo simulation

The Monte Carlo simulation [Rubinstein \(1981\)](#) is the most robust of all the presented methods. Samples are generated with respect to the probability density function of each variable and for each sample the response of the structure is determined. An unbiased estimator  $\tilde{p}_f$  of the failure probability is given by

$$\tilde{p}_f = \frac{1}{N} \sum_{i=1}^N I(\mathbf{y}), \quad (4)$$

where  $N$  is the number of samples. The expected value and the variance of the estimator are given by

$$E[\tilde{p}_f] = p_f \quad (5)$$

$$Var[\tilde{p}_f] = \frac{1}{N} Var[I(\mathbf{x})] \quad (6)$$

$$Var[I(\mathbf{x})] = E[(I(\mathbf{x}))^2] - (E[I(\mathbf{x})])^2. \quad (7)$$

The variance  $Var[\tilde{p}_f]$  and the statistical error  $e$

$$e = \frac{Var[\tilde{p}_f]}{E[\tilde{p}_f]} = \frac{1}{\sqrt{N}} \frac{\sqrt{Var[I(\mathbf{x})]}}{p_f} \quad (8)$$

are an indicator for the quality of the estimator  $\tilde{p}_f$ . Especially for small failure probabilities  $p_f$  a large number of samples is necessary in order to obtain a good estimator. Assuming that e.g.  $Var[I(\mathbf{x})] = p_f = 10^{-4}$  a total number of  $N = 4 \cdot 10^4$  samples is required to obtain a statistical error  $e = 0.5$  and  $N = 1 \cdot 10^6$  samples for a statistical error  $e = 0.1$ . For complex problems with a computational expensive algorithm for the calculation of a single sample the evaluation of such a large number of samples is not acceptable.

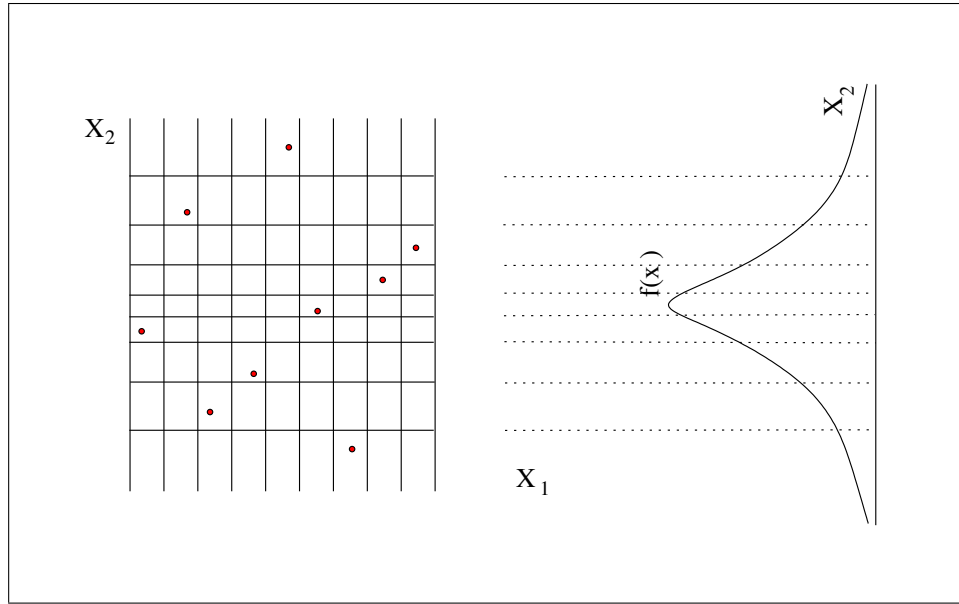


Figure 1: Latin hypercube

### 1.3 Latin Hypercube Sampling

Latin Hypercube sampling was introduced by [McKay u. a. \(1979\)](#). The domain of each random variable is decomposed into intervals with equal probability. The number of intervals corresponds to the number of samples. One value from each interval is selected at random with respect to the probability density in the interval. Combining the intervals of a random variables, the so called hypercubes are formed. This is illustrated in [fig.1](#) and in a similar way the representative point within one hypercube. The  $n$  representantive values obtained for  $X_1$  are paired in a random manner (equally likely combinations) with the  $n$  values of  $X_2$  and so on until  $X_n$ . These  $n$  random combinations are called the Latin Hypercube Samples. If the samples are combined in a random manner, artificial correlations are introduced, which can be avoid by regrouping the samples [Iman und Conover \(1982\)](#); [Stein \(1987\)](#). An estimate of the failure probability can be calculated (analoge to plain Monte Carlo) by

$$\tilde{p}_f = \frac{1}{N} \sum_{i=1}^N I(\mathbf{y}), \quad (9)$$

### 1.4 Adaptive Sampling

This procedure [Bucher \(1988\)](#) is based on the Importance sampling technique, which concentrates the samples within the regions close to the failure domain. The integral in [eq.\(3\)](#) is written as

$$p_f = \int_{\Omega} I(\mathbf{x}) f(\mathbf{x}) d\mathbf{x} \quad (10)$$

$$= \int_{\Omega} \left[ I(\mathbf{y}) \frac{f(\mathbf{x})}{h(\mathbf{x})} \right] h(\mathbf{x}) d\mathbf{x} \quad (11)$$

$$= \int_{\Omega} \bar{I}(\mathbf{y}) h(\mathbf{y}) d\mathbf{y}, \quad (12)$$

with  $\bar{I}$  a weighted indicator function and  $h(\mathbf{y})$  the joint probability density function of the random variables  $\mathbf{Y}$ . An estimator of the failure probability  $p_f$  is, analogue to the Monte Carlo simulation, obtained by generating samples with respect to the joint probability density function  $h(\mathbf{y})$ . An unbiased estimator for the failure probability  $p_f$  is given by

$$\tilde{p}_f = \frac{1}{N} \sum_{i=1}^N I(\mathbf{y}) \frac{f(\mathbf{y})}{h(\mathbf{y})} \quad (13)$$

with

$$E[\tilde{p}_f] = p_f \quad (14)$$

$$Var[\tilde{p}_f] = \frac{1}{N} \int_{\Omega} I^2(\mathbf{y}) \frac{f^2(\mathbf{y})}{h(\mathbf{y})} d\mathbf{y} - \frac{p_f^2}{N} \quad (15)$$

The expected value  $E[\tilde{p}_f]$  is independant of the sampling density function  $h(\mathbf{y})$ , whereas the variance of the estimator strongly depend on  $h(\mathbf{y})$ . The optimal sampling density function  $h(\mathbf{y})$ , for which  $Var[\tilde{p}_f] = 0$ , is given by

$$h(\mathbf{y}) = \frac{1}{p_f} I(\mathbf{y}) f(\mathbf{y}) \quad (16)$$

Since the failure probability  $p_f$  is a priori the unknown variable, the optimal sampling density function  $h(\mathbf{y})$  is approximated as follows. The sampling density function for the first iteration step is either the original density function  $f(\mathbf{y})$  or an enlarged density function, where the variance of the original function is multiplied by a factor between 1 and 3 in order to have more samples within the failure domain. The latter approach is especially suitable for small failure probabilities. After the first iteration the new sampling density function  $h(\mathbf{y})$  is determined, so that the first and second moments coincide with the ones of the samples, that are within the failure domain.

$$E[\mathbf{Y}] = E[\mathbf{X} | g(\mathbf{x}) \leq 0] \quad (17)$$

$$E[\mathbf{Y}\mathbf{Y}^T] = E[\mathbf{X}\mathbf{X}^T | g(\mathbf{x}) \leq 0] \quad (18)$$

This iteration process is repeated several times. In practise 3 iterations are in general sufficient.

The problem of this method is a robust estimation of the covariance matrix in eq.(18). If the number of samples within the failure domain is relatively small, which can be either due to a small number of total samples or a small failure probability, the covariance matrix is often not positive definite, and as a result the transformation of  $\mathbf{Y}$  to uncorrelated random variables is not possible.

In order to avoid this problem, a restriction for the number of samples  $N$  is introduced:

$$n < p_f \cdot N, \quad (19)$$

where  $p_f$  is the expected failure probability and  $n$  the number of random variables. This restriction means, that at least  $n$  samples within the failure domain are expected and that this samples can be used for the calculation of the covariance matrix of the adaptive joint probability density function. Obviously an estimation of the failure probability is a priori required, which might in many cases not be given. An approximation can either be obtained by engineering experience, by other methods as e.g. FORM or by using the result of the initial Monte Carlo simulation. For any configuration that violates this condition no failure probability is calculated. Furthermore the number of samples within the failure domain  $N_f$  for each iteration step is required to be larger than twice the number of random

variables,

$$N_f > 2 \cdot n. \quad (20)$$

Otherwise no further iteration is performed and the failure probability is calculated from the current iteration step.

## 1.5 FORM - First Order Reliability Method

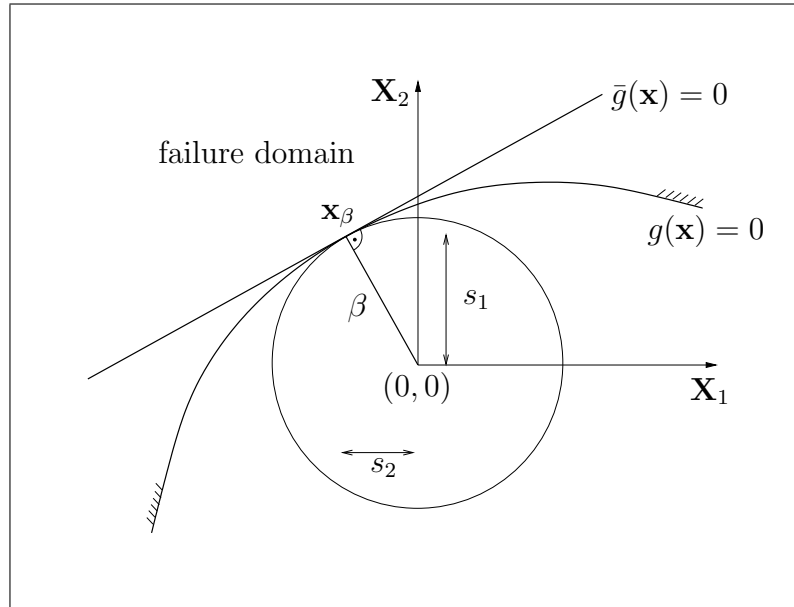
FORM [Shinozuka \(1983\)](#); [Hasofer und Lind \(1974\)](#); [Hohenbichler und Rackwitz \(1988\)](#); [Tvedt \(1983\)](#); [Breitung \(1991\)](#) is one of the most effective methods in reliability analysis, since the number of required function evaluations is relatively small. In order to determine the failure probability  $p_f$  in eq.(1) the limit state function  $g(\mathbf{x})$  can be approximated by a Taylor expansion. The optimal expansion point is the design point  $\mathbf{x}_\beta$ , which is the point on  $g(\mathbf{x})$  closest to the origin. The distance  $\beta$  to the origin is called the reliability index.

The hyperplane  $g(\mathbf{x}) = 0$  is characterized by

$$g(\mathbf{x}) = -\sum_{i=1}^N \frac{x_i}{s_i} + 1 = 0 \quad (21)$$

Furthermore the normal  $\vec{b}$  to the hyperplane is given by  $\left[-\frac{1}{s_1}, -\frac{1}{s_2}, \dots, -\frac{1}{s_n}\right]^T$ , which finally gives the distance  $\beta$  to the origin:

$$\beta = \frac{1}{\sqrt{\sum_{i=1}^n \frac{1}{s_i^2}}} \quad \Rightarrow \quad \sum_{i=1}^n \frac{1}{s_i^2} = \frac{1}{\beta^2} \quad (22)$$



**Figure 2:** approximation of the limit state function  $g(\mathbf{x})$  by a linear function  $\bar{g}(\mathbf{x})$  in standard Gaussian space

The random variable  $Z = g(\mathbf{X})$  is called the safety margin and its linear approximation  $\bar{Z}$  is given by

$$\bar{Z} = \bar{g}(\mathbf{X}) = - \sum_{i=1}^N \frac{x_i}{s_i} + 1. \quad (23)$$

It is normal distributed, since it is the sum of independant normal distributed variables.

$$F_{\bar{Z}}(\bar{z}) = \Phi \left( \frac{\bar{z} - E[\bar{Z}]}{\sigma_{\bar{Z}}} \right) \quad (24)$$

The statistical moments are

$$E[\bar{Z}] = 1 \quad (25)$$

$$\begin{aligned} \sigma_{\bar{Z}}^2 &= E[(\bar{Z} - E[\bar{Z}])^2] = E \left[ \left( \sum_{i=1}^N \frac{X_i}{s_i} \right)^2 \right] \\ &= E \left[ \sum_{i=1}^N \sum_{j=1}^N \frac{X_i X_j}{s_i s_j} \right] = E \left[ \sum_{i=1}^N \frac{X_i^2}{s_i^2} \right] = E \left[ \sum_{i=1}^N \frac{1}{s_i^2} \right] \\ &= \frac{1}{\beta^2} \end{aligned} \quad (26)$$

The failure probability is given by

$$p_f = P[Z \leq 0] \approx P[\bar{Z} \leq 0] \quad (27)$$

$$= \int_{-\infty}^0 f_{\bar{Z}}(\bar{z}) d\bar{z} = F_{\bar{Z}}(0) \quad (28)$$

$$= \Phi \left( \frac{-E[\bar{Z}]}{\sigma_{\bar{Z}}} \right) = \Phi \left( \frac{1}{\frac{1}{\beta}} \right) \quad (29)$$

$$= \Phi \left( \frac{1}{\beta} \right) \quad (30)$$

In general FORM gives a good approximation of the failure probability. The first problem of this method is the determination of the design point. In SLang the routine NLPQL is used, which is based on a sequential quadratic programming (SQP) method. For further details see [Schittkowski \(1985\)](#). In certain cases (see examples **reference examples here**) the determined design point deviates considerably from the theoretical design point. A second problem is the existance of multiple design points, which can not be handled by this method. Furthermore the probability density function is approximated by a linear function, which results in additional errors for non linear functions.

## 1.6 Importance Sampling Procedure Using Design points (ISPUD)

Another possibility to determine the sampling density function within an Importance sampling approach is used in the ISPUD algorithm [Bourgund und Bucher \(1986\)](#). At first the design point is determined, which is analog to the FORM performed using NLPQL. The mean value of the normal distributed sampling density function is the design point, whereas the variance is the same as the original density function. Although the method is applicable in Gaussian as well as in original space, it seems to be advantageous to perform

the sampling in Gaussian space, especially if the random variables are correlated or not normal distributed. In this case the variance of the sampling density function (which equals 1 for the random variables transformed to Gaussian space) can be increased by a certain factor in order to spread the samples further from the design point. The failure probability is calculated using eq.(13). The advantage of this method is its independance with respect to the non linearity of the limit state function close to the design point, but similar to FORM the determination of the correct design point or the existence of multiple design points can cause errors.

## 1.7 Directional Sampling

For the directional sampling Bjerager (1988); Ditlevsen und Bjerager (1989) the random variables are transformed to uncorrelated variables in Gaussian space. Starting from the origin, which corresponds to the mean value, random direction vectors are created. In these directions the point of failure is determined by using a bisection algorithm.

Each point  $\mathbf{x}$  is written as  $\mathbf{x} = r\mathbf{a}$ , where  $r$  is the distance from the origin and  $\mathbf{a}$  is a unit direction vector. The vector  $\mathbf{a}$  is described by  $n-1$  independant variables  $\Phi_i$ , which are the  $n$ -dimensional spherical coordinates.

### Probability density function

The joint density function of the variables  $\mathbf{X}$  in standard Gaussian space is given by

$$f_{\mathbf{X}}(\mathbf{x}) = \frac{1}{(2\pi)^{\frac{n}{2}}} \mathbf{e}^{(-\frac{1}{2}\mathbf{x}^T\mathbf{x})}. \quad (31)$$

Transformation of the variables to  $n$ -dimensional spherical coordinates yields

$$f_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} = f_{\mathbf{X}}(\mathbf{x}(r, \phi)) r^{n-1} dr d\phi \quad (32)$$

$$= f_{R,\Phi}(r, \phi) dr d\phi \quad (33)$$

$f_{\mathbf{X}}(\mathbf{x}(r, \phi)) r^{n-1} = f_{R,\Phi}(r, \Phi)$  is rotationally symmetric, i.e. independent of  $\phi$ , since

$$\mathbf{x}^T\mathbf{x} = (r\mathbf{a})^T (r\mathbf{a}) = r^2\mathbf{a}^T\mathbf{a} = r^2. \quad (34)$$

As a result the random variables  $R$  and  $\Phi$  are independant.

$$f_{R|\Phi}(r|\phi) = f_R(r) \quad (35)$$

$$f_{R,\Phi}(r, \phi) = f_{R|\Phi}(r|\phi)f_{\Phi}(\phi) = f_R(r)f_{\Phi}(\phi) \quad (36)$$

Due to rotational symmetry of  $f_{R,\Phi}$ ,  $f_{\Phi}$  must have identical values for any  $\phi$ . As a result the probability density function is constant and its value is the inverse of the surface area  $S_n$  of the  $n$ -dimensional unit hypersphere.

$$f_{\Phi}(\phi) = \frac{1}{S_n} = \frac{\Gamma(\frac{n}{2})}{2\pi^{\frac{n}{2}}} \quad (37)$$

Using eq.(32,33),(36) and (31) it follows, that:

$$f_R(r) = \frac{f_{R,\Phi}(r, \phi)}{f_{\Phi}(\phi)} \quad (38)$$

$$= \frac{S_n r^{n-1} \mathbf{e}^{(-\frac{r^2}{2})}}{(2\pi)^{\frac{n}{2}}} \quad (39)$$



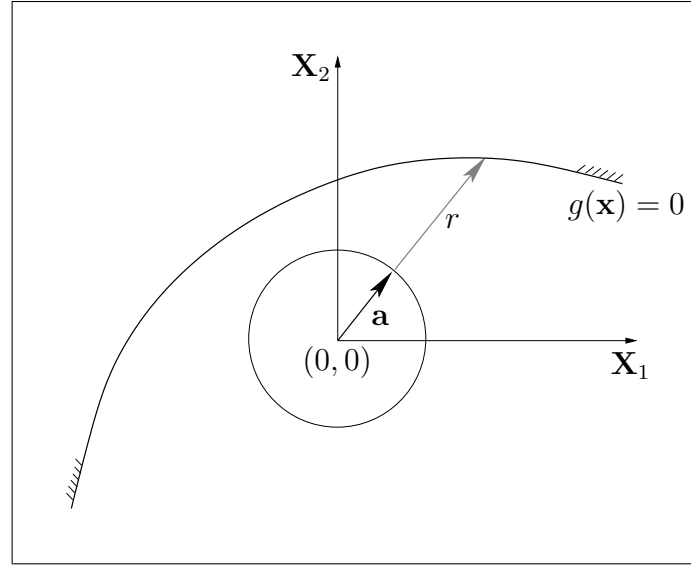


Figure 3: directional sampling

### Probability of failure

The conditional failure probability  $p_f(|\Phi = \phi)$  for a given direction  $\phi$  can be determined analytically

$$\begin{aligned}
 p_f(|\Phi = \phi) &= \int_{R^*(\phi)}^{\infty} f_{R,\Phi}(r, \phi | \Phi = \phi) dr \\
 &= \int_{R^*(\phi)}^{\infty} f_R(r) dr \\
 &= \int_{R^*(\phi)}^{\infty} \frac{S_n r^{n-1} e^{-\frac{r^2}{2}}}{(2\pi)^{\frac{n}{2}}} dr \\
 &= 1 - \chi_n^2(R^*(\phi)^2),
 \end{aligned} \tag{40}$$

where  $\chi_n$  is the cumulative Chi-Square-Distribution function.

Assuming, that the origin is not inside the failure domain and that there is a unique point of failure for any direction  $\mathbf{a}$ , the failure probability  $p_f$  can be expressed as

$$p_f = \int_{\Omega}^n \int I(\mathbf{x}) f_{\mathbf{x}}(\mathbf{x}) d\mathbf{x} \tag{41}$$

$$= \int_{S_n}^{n-1} \int \int_{r=0}^{\infty} I(r, \phi) f_{R,\Phi}(r, \phi) dr d\phi \tag{42}$$

$$= \int_{S_n}^{n-1} \int \int_{r=R^*}^{\infty} f_R(r) f_{\Phi}(\phi) dr d\phi \tag{43}$$

$$= \int_{S_n}^{n-1} \int p_f(|\Phi = \phi) f_{\Phi}(\phi) d\phi. \tag{44}$$

An unbiased estimator of the failure probability is given by

$$\tilde{p}_f = \frac{1}{N} \sum_{j=1}^N p_f(|\Phi = \phi) \quad (45)$$

$$= \frac{1}{N} \sum_{j=1}^N [1 - \chi_n^2(R^*(\phi_j)^2)], \quad (46)$$

where the sample directions  $\phi_j$  are generated by simulating samples  $\mathbf{x}_j$  according to the density functions in standard Gaussian space and then calculating the direction vector as  $\mathbf{a} = \frac{\mathbf{x}_j}{\|\mathbf{x}_j\|}$ , from which the angles  $\phi_j$  can be obtain. In practice, only the direction vector  $\mathbf{a}$  is determined and the point of failure in this direction ( $R^*(\mathbf{a})\mathbf{a}$ ) is determined by a bisection algorithm. The one-dimensional integral in eq.(46) can be evaluated by using the formular of [Abramowitsch und Stegun \(1972\)](#)

$$1 - \chi_n^2(R^{*2}) = 1 - \chi_{n-2}^2(R^{*2}) + \frac{\left(\frac{R^{*2}}{2}\right)^{n/2-1} e^{-r^2/2}}{\Gamma\left(\frac{n}{2}\right)}, \quad n > 2. \quad (47)$$

For  $n = 1$  an exponential function and for  $n = 2$  the cumulative standard Gaussian distribution has to be evaluated, which is done numerically.

## 1.8 Adaptive Directional Sampling

The adaptive directional sampling is an importance sampling method (see [1.4](#)). After an inital directional sampling step the simulation density function is adapted the way, that the number of samples in the direction of the failure domain from the previous iteration step is increased. The adaptation of the simulation density function for the second (and further steps) is performed by using the points of the previous step. The first two statistical moments of these points is equal to the parameters of the new simulation density function  $h_{\Phi}(\phi)$ .

$$p_f = \int_{S_n}^{n-1} \int_{r=R^*}^{\infty} f_R(r) f_{\Phi}(\phi) dr d\phi \quad (48)$$

$$= \int_{S_n}^{n-1} \int_{r=R^*}^{\infty} f_R(r) \left[ \frac{f_{\Phi}(\phi)}{h_{\Phi}(\phi)} \right] h_{\Phi}(\phi) dr d\phi \quad (49)$$

$$= \int_{S_n}^{n-1} p_f(|\Phi = \phi) \left[ \frac{f_{\Phi}(\phi)}{h_{\Phi}(\phi)} \right] h_{\Phi}(\phi) d\phi \quad (50)$$

An unbiased estimator of the failure probability is given by

$$\tilde{p}_f = \frac{1}{N} \sum_{j=1}^N \left\{ p_f(|\Phi = \phi_j) \left[ \frac{f_{\Phi}(\phi_j)}{h_{\Phi}(\phi_j)} \right] \right\} \quad (51)$$

$$= \frac{1}{N} \sum_{j=1}^N \left\{ [1 - \chi_n^2(R^*(\phi_j)^2)] \left[ \frac{f_{\Phi}(\phi_j)}{h_{\Phi}(\phi_j)} \right] \right\}. \quad (52)$$

For small failure probabilities it seems to be advantageous to use already for the initial simulation an enlarged simulation density function in order to ensure, that a sufficient number of points is within the failure domain. Otherwise the estimation of the covariance matrix is impossible.

## 1.9 Response surface methods

The idea of the response surface methods is, that the complex response of a (mechanical) system can be approximated by a simple function of the input variables  $X_i$ . Traditionally a polynomial approximation is used

$$g(\mathbf{x}) = \theta_0 + \sum_{i=1}^n \theta_i x_i + \sum_{j=1}^n \sum_{i=k}^n \theta_{jk} x_j x_k + \dots + \epsilon \quad (53)$$

, where  $\epsilon$  is an approximation error and  $\theta$  are the coefficients, which are determined by evaluating the response  $g(x_i)$  for certain points  $x_i$ . A linear system of equations is obtained, which can be solved e.g. with a least squares approach. Different approaches can be used to determine the location of the points used for the approximation of the limit state function, e.g. by

$$\mathbf{x}_{2i} = (\mu_1, \mu_2, \dots, \mu_i + \sigma_i, \dots, \mu_n) \quad (54)$$

$$\mathbf{x}_{2i+1} = (\mu_1, \mu_2, \dots, \mu_i - \sigma_i, \dots, \mu_n). \quad (55)$$

It is often sufficient to decide, if a point is located within the safe ( $g(\mathbf{x}) \geq 0$ ) or in the failure domain ( $g(\mathbf{x}) < 0$ ). Applying the concept of response surfaces, the limit state function  $g(\mathbf{x})$  can be approximated by

$$g(\mathbf{x}) = \theta_0 + \sum_{i=1}^n \theta_i x_i + \sum_{j=1}^n \sum_{i=k}^n \theta_{jk} x_j x_k + \dots + \epsilon = 0 \quad (56)$$

The traditional approximation of response surfaces with polynomials has the disadvantage, that if the number of sampling points is higher than the number of coefficients  $\theta$ , the response surface does not pass through the points itself.

### 1.9.1 Approximation of the limit state function by tangential hyperplanes

Another possible approximation of the response surface is the use of hyper planes. Assuming, that a point  $\mathbf{p}$  is on the failure surface. The normal  $\mathbf{n}$  of the corresponding hyperplane is described by the vector from a point  $\mathbf{m}$  to  $\mathbf{p}$  and can be written in its Hessian form as

$$(\mathbf{p} - \mathbf{x}) \mathbf{n} = 0 \quad (57)$$

An illustration of the response surface is given in figure 4. The point  $\mathbf{m}$  has to be inside the safe domain and can e.g. be set to the mean value of the random variables. Given a set of points  $\mathbf{p}_i$  on the limit state function the test, whether a point  $\mathbf{x}$  is inside the safe domain is performed as follows.

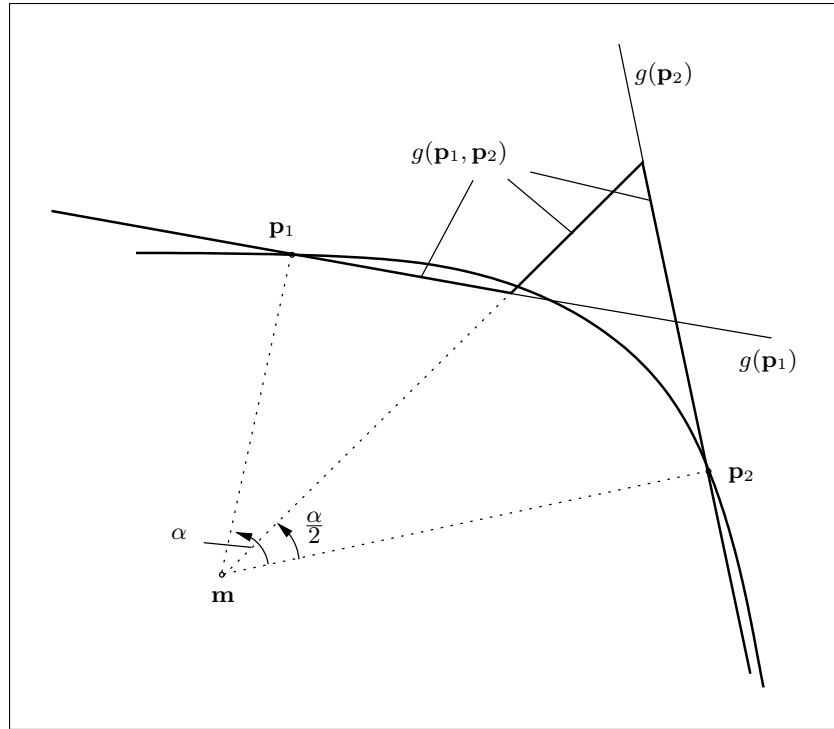
- determine the closest point  $\mathbf{p}_j$  on the response surface by

$$\mathbf{p}_j = \left( \mathbf{p}_i : \cos(\phi_i) = \frac{\mathbf{p}_i \mathbf{x}}{\|\mathbf{p}_i\| \|\mathbf{x}\|} \longrightarrow \max \right) \quad (58)$$

- determine the indicator function  $I$  by evaluation of eq.(57).

$$I(\mathbf{x}) = \begin{cases} 1 & : (\mathbf{p} - \mathbf{x}) \mathbf{n} \leq 0 \\ 0 & : (\mathbf{p} - \mathbf{x}) \mathbf{n} > 0 \end{cases} \quad (59)$$

to determine, if a point is inside ( $I=1$ ) or outside ( $I=0$ ) the failure domain.

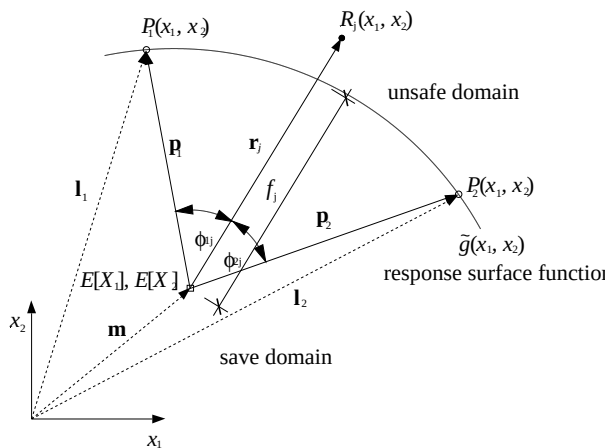


**Figure 4:** response surface approximated by hyperplanes normal to the points  $\mathbf{p}_i$  on limit state function

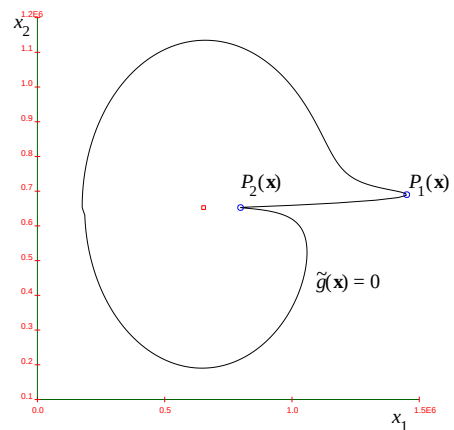
### 1.9.2 Approximation of the limit state function using weighted radii

The approximation of the limit state function using weighted radii is based on weighting the direction vector from a point  $\mathbf{m}$  to a point  $\mathbf{x}$  using the points  $\mathbf{p}_i$  on the response surface. The weight is determined as follows. For a point  $\mathbf{x}$  the angles  $\phi_i$  between the vectors  $(\mathbf{p}_i - \mathbf{m})$  and  $(\mathbf{x} - \mathbf{m})$  are given by

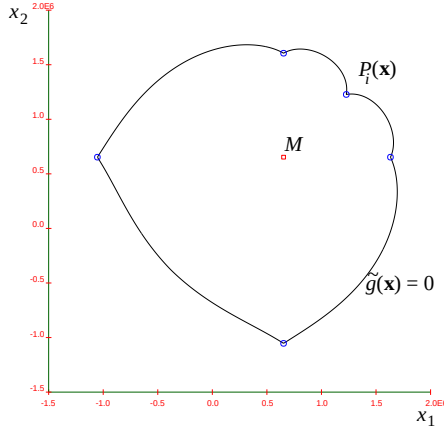
$$\cos(\phi_i) = \frac{(\mathbf{p}_i - \mathbf{m})(\mathbf{x} - \mathbf{m})}{\|\mathbf{p}_i - \mathbf{m}\| \|\mathbf{x} - \mathbf{m}\|} \quad \text{with } 0 \leq \phi_i < \pi \quad (60)$$



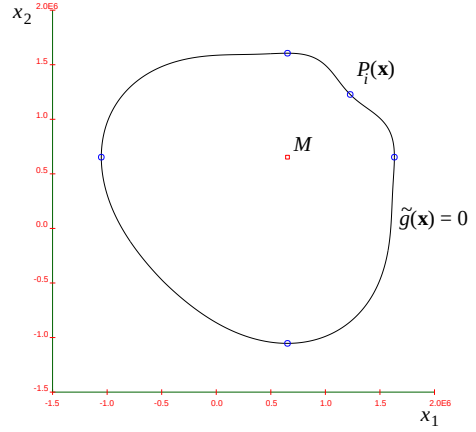
**Figure 5:** Limit state interpolation using weighted radii.



**Figure 6:** Interpolation without difficulties in near of discontinuities.



**Figure 7:** Response surface by using linear weights.



**Figure 8:** Response surface by using non-linear weights.

One possibility to determine the weight  $w_i(\mathbf{x})$  and the scaling factor  $f(\mathbf{x})$  is

$$w_i = \frac{1}{\phi_i} \quad (61)$$

$$(62)$$

$$f = \frac{\sum_i \|\mathbf{p}_i\| w_i}{\sum_i w_i} \quad (63)$$

Finally the limit state function is approximated by

$$g(\mathbf{x}) = \mathbf{m} + f(\mathbf{x}) \frac{\mathbf{x} - \mathbf{m}}{\|\mathbf{x} - \mathbf{m}\|} = 0 \quad (64)$$

If only one supporting point  $\mathbf{p}$  is present, the response surface is a hypersphere. The weighting function has a pole for  $\phi_i = 0$ . By introducing a small parameter  $\epsilon$  in the weighting function this singularity can be removed.

$$w_i = \frac{1}{\phi_i + \epsilon} \quad (65)$$

The disadvantage is, that the interpolation does not pass through the points of support, but by setting  $\epsilon$  sufficiently small, this error is negligible.

Instead of using all points of support for the calculation of the weighting factor  $f$ , it is often sufficient to use only the  $m$  points with the highest weighting function. In general  $m = n..2n$ , where  $n$  is the number of random variables, results in a good approximation. It is to be noted, that the response surface should preferential be computed and analyzed in standard Gaussian space.

The response surface function for some limit state check points using linear weights is shown in Fig. 7. Introducing nonlinear weights such as

$$w_{ij} = \left( \frac{1}{\phi_{ij} + \epsilon} \right)^2$$

will form a differentiable function in the supporting points, as shown in Fig. 8. Using the presented interpolation, there are no difficulties in near of discontinuities, as shown in Fig. 6, in contrast to polyhedral and polynomial response surface functions.

### 1.9.3 Calculating the points of support for the limit state function

The determination of the points of support used for the approximation of the limit state function can be performed with a deterministic, a stochastic or a combined approach. In the deterministic approach direction vectors  $d_i$  from the mean value (or the origin in standard Gaussian space) are created according to a user defined scheme, e.g. by searching along the axis from the mean value, which results in  $2n$  search directions ( $d_1..d_4$  for a 2D-example) or by searching along the direction to the center of all quadrants ( $d_5..d_8$ ), which results in  $2^n$  search directions.

$$\begin{aligned} d_1 &= (1, 0) & d_2 &= (-1, 0) & d_3 &= (0, 1) & d_4 &= (0, -1) \\ d_5 &= (1, 1) & d_6 &= (-1, 1) & d_7 &= (-1, -1) & d_8 &= (1, -1) \end{aligned} \quad (66)$$

Obviously the number of directions increases dramatically as the number of random variables increases and for  $n = 20$  already  $\approx 10^6$  search directions have to be considered. Adaptive schemes are required for higher dimensions, where e.g. after the search along the axis only the quadrants are considered, that are enclosed by failure points on the axis with a high probability density function.

In a stochastic approach random direction vectors are calculated according to a probability density function. Obviously the number of required directions in order to obtain a good correspondance between the exact and approximated limit state function is also large for a large number of random variables. In this investigation a combined approach has been followed. Starting from the deterministic approach and additional stochastic direction vectors with a uniform distribution over the hypersphere an adapted probability density function (analog to adaptive directional sampling in section 1.8) is calculated and in the second step the direction vectors are calculated with respect to this adapted probability density function. The approximate location of the critical regions is thereby determined by the deterministic approach plus additional stochastic directions, and the adaptation increases the number of samples close to the design point.

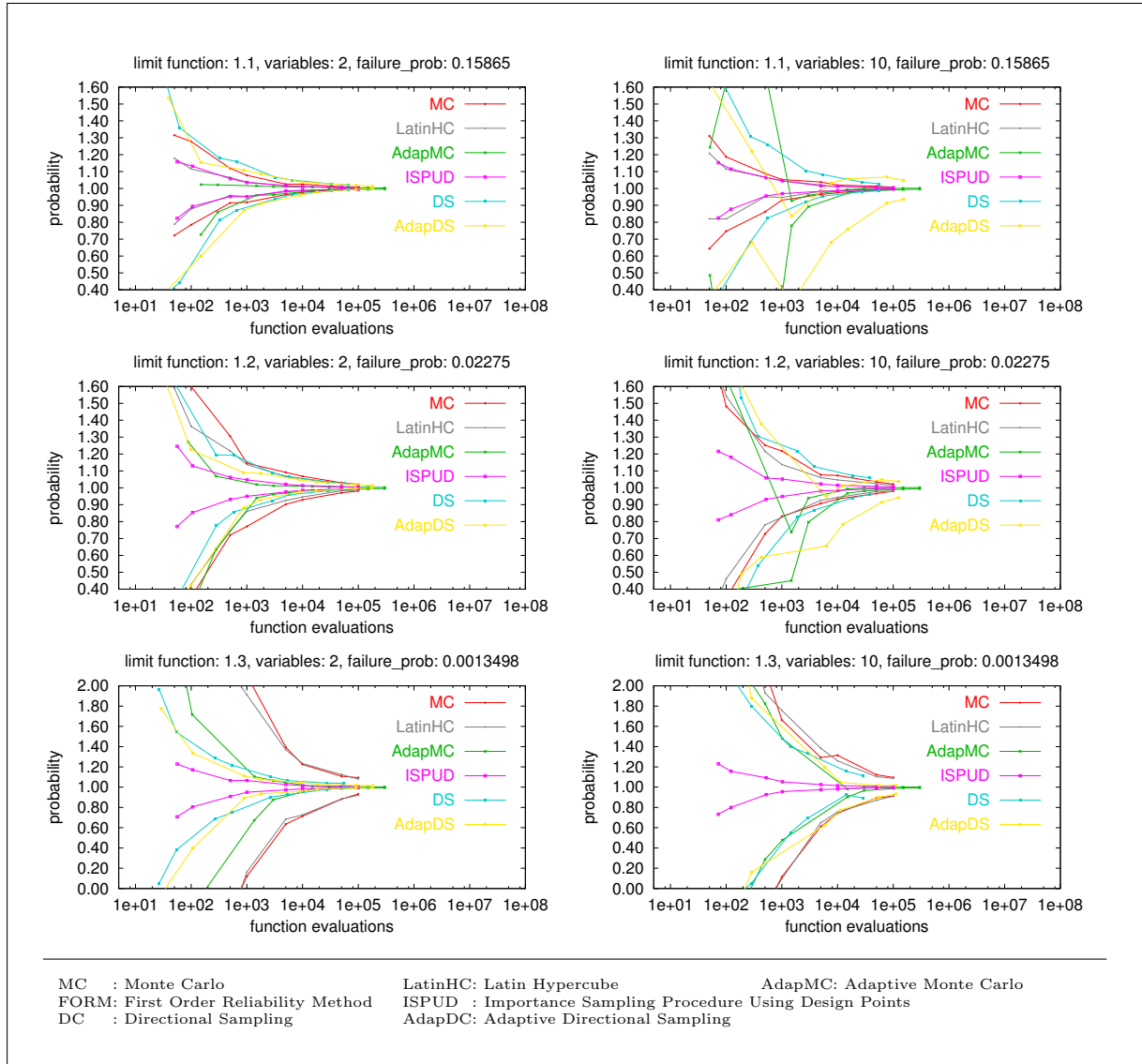
## 2 Limit state functions

In most practical applications the characteristics of the limit state function as e.g. its shape are not known a priori. Furthermore the limit state function is often not an analytic function but derived from a numerical model. It is therefore necessary to investigate the performance of different methods with respect to the following criteria

- probability level
- number of random variables
- multiple  $\beta$ -points
- curvature of the limit state function
- applicability in standard/original space
- noisy limit state function

### 2.1 Limit state function 1

[bt] The first limit state function is used to compare the performance of the methods



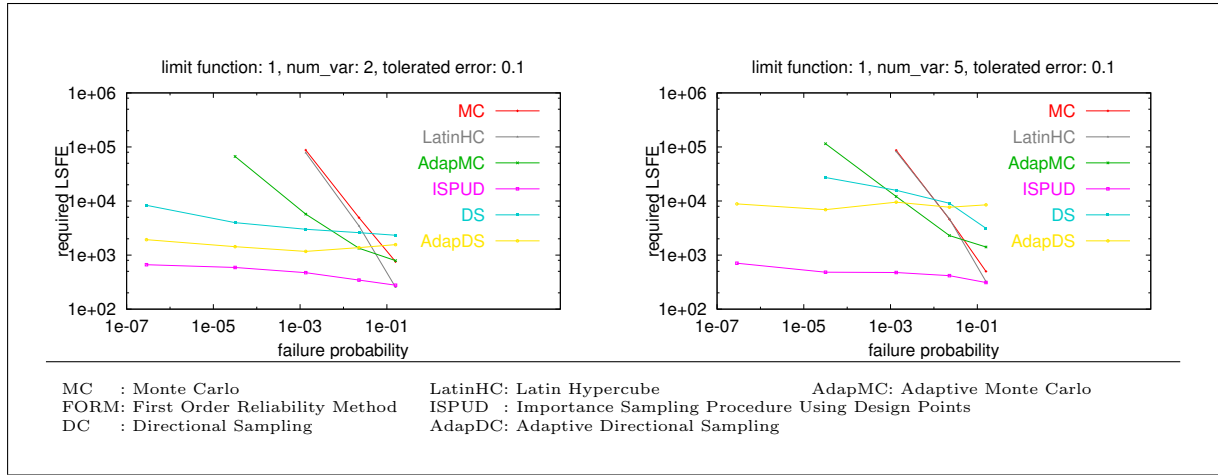
**Figure 9:** intervall of the obtained failure probability  $\left[ \frac{\mu \pm \sigma}{\mu_{exact}} \right]$  for limit function 1 with 2 or 10 random variables and different failure probabilities

with respect to different failure probabilities and different numbers  $n$  of random variables. The influence of other parameters is reduced by using normal distributed variables and a linear limit state function (a hyperplane in higher dimensions).

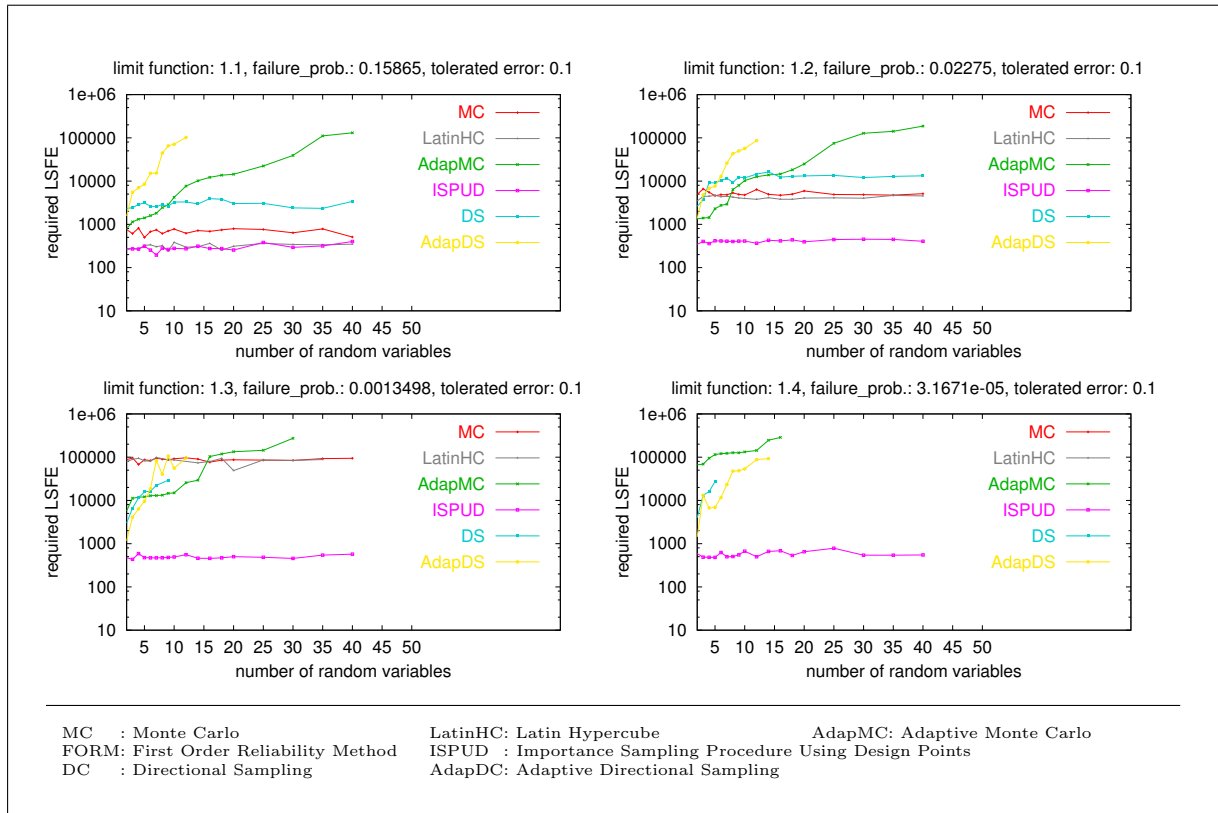
$$g_1 = \beta\sqrt{n} - \sum_{i=1}^n U_i = 0 \quad (67)$$

$U_i$  : independent normal distributed  $\mu = 0, \sigma^2 = 1$

In Fig.9 the calculated failure probability for different methods is illustrated. The two graphs for each method indicate the intervall  $[\mu \pm \sigma]$ . The parameters are obtained by performing 100 experiments with the same conditions (approximate number of limit state function evaluations, number of random variables, failure probability). From these 100 experiments the mean and the standard deviation are calculated. Obviously with a number of 100 experiments only an approximation of the accuracy of the estimator can be determined. For this limit state function the Design Point could be determined exactly using NLPQL, and due to the linearity of the limit state function the result obtained by



**Figure 10:** required number of samples for  $tol = 0.1$  with 2 and 5 random variables for varying probability level

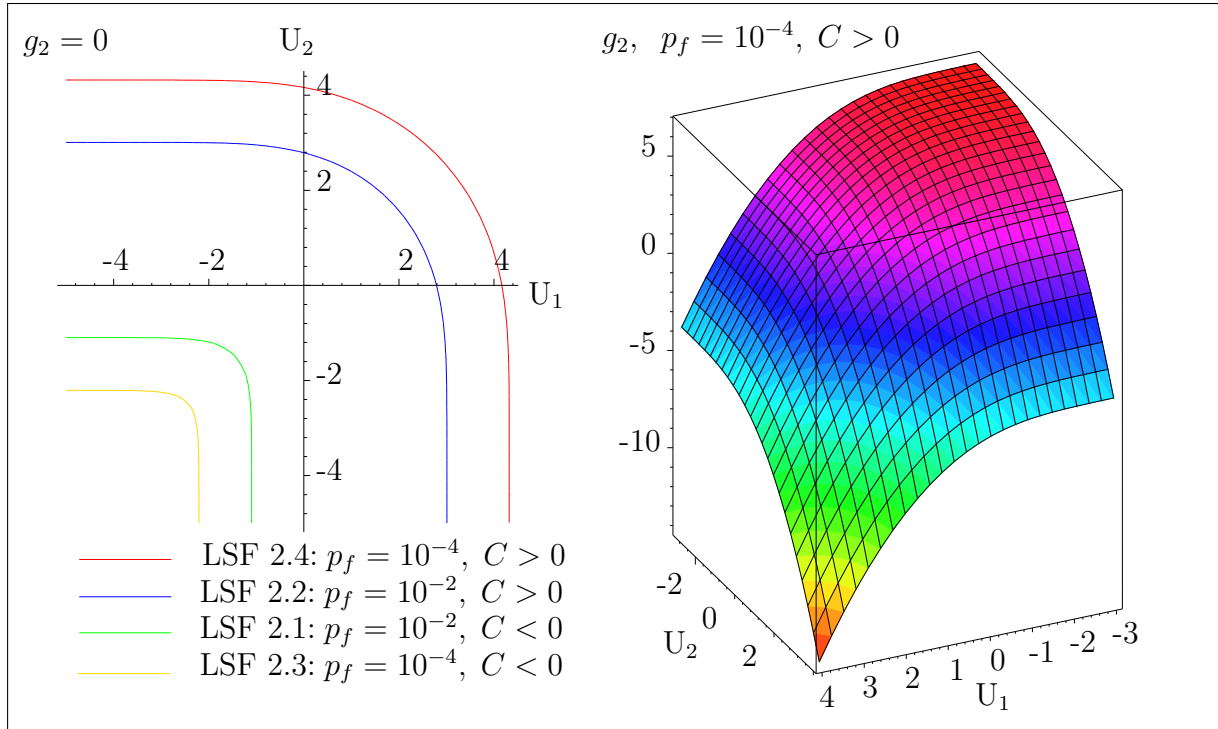


**Figure 11:** required number of samples for  $tol=0.1$  for varying dimension (number of random variables)

FORM is exact. Furthermore it is observed (as predicted theoretically), that the accuracy of the plain Monte Carlo simulation is independent of the dimension (number of random variables), but strongly depends on the failure probability. This is further illustrated in Fig. 10, where for a given tolerance  $tol = 0.1$  the number of required samples is plotted as a function of the probability level, so that

$$1 - tol > \frac{\mu - \sigma}{\mu_{exact}} < \frac{\mu + \sigma}{\mu_{exact}} < 1 + tol. \quad (68)$$





**Figure 12:** limit state function  $g_2$  for varying level of probability and positive and negative curvature

The Importance Sampling Procedure Using Design Points gives accurate results, since the Design Point can be determined exactly. Furthermore the points on the failure surface with a relatively high probability are all close to this designpoint. Adaptive Sampling gives accurate results, if in the first iteration step a sufficient number of points is within the failure domain and as a result a better estimate of the correlation between the variables is determined. Latin Hypercube is in general slightly better than plain Monte Carlo, but the general disadvantage (strong dependence on the probability level) is observed. Directional Sampling is for this limit state function (hyperplane)independent of the probability level. This is due to the fact, that by modifying the distance of the hyperplane to the origin (which corresponds to different levels of probability) only the number of iterations, that is needed for a random sample direction is influenced. As a result for the same number of directions slightly more iterations are necessary to find the limit state function.

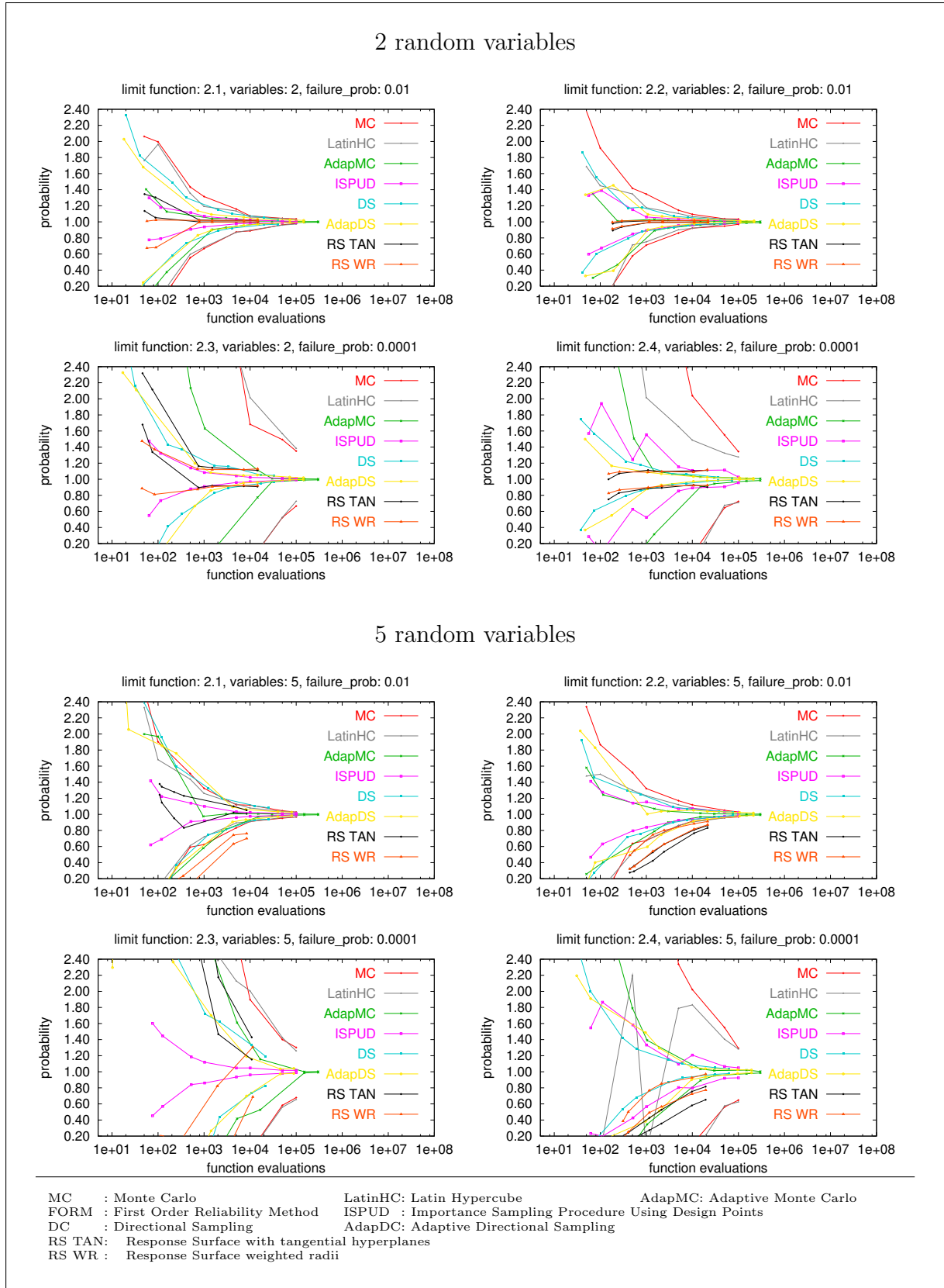
## 2.2 Limit state function 2

The second limit state function is linear (LSF) in the original space, but a transformation to standard Gaussian space results in a nonlinear function.

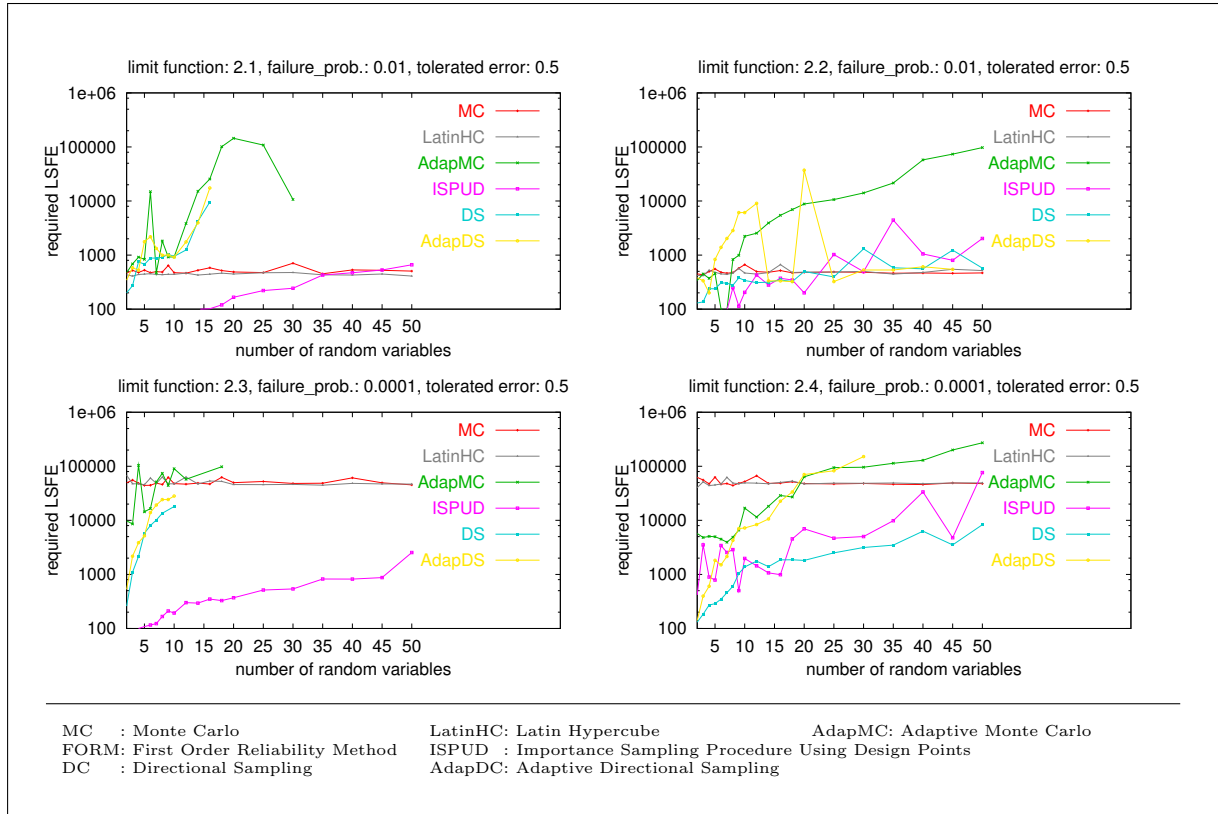
$$g_2 = \pm C \mp \sum_{i=1}^n X_i = 0 \quad (69)$$

$X_i$  : independent exponential distributed  $\lambda = 1$

As illustrated in Fig.12 the LSF is for  $C > 0$  convex with respect to the origin and concave for  $C < 0$ . The relative mean probability with the corresponding intervall  $\left[ \frac{\mu \pm \sigma}{\mu_{exact}} \right]$  is illustrated in fig.13. For this limit state function the performance of the response surface methodes has been additionally investigated. For 2 random variables with a high failure probability the response surface methods a superior to all other methods. This is on the



**Figure 13:** intervall of the obtained failure probability  $\left[ \frac{\mu \pm \sigma}{\mu_{exact}} \right]$  for limit function 2 with 2 and 5 random variables, failure probabilities of 0.01 (LSF 2.1 and 2.2) and 0.0001 (LSF 2.3 and 2.4) and different curvatures



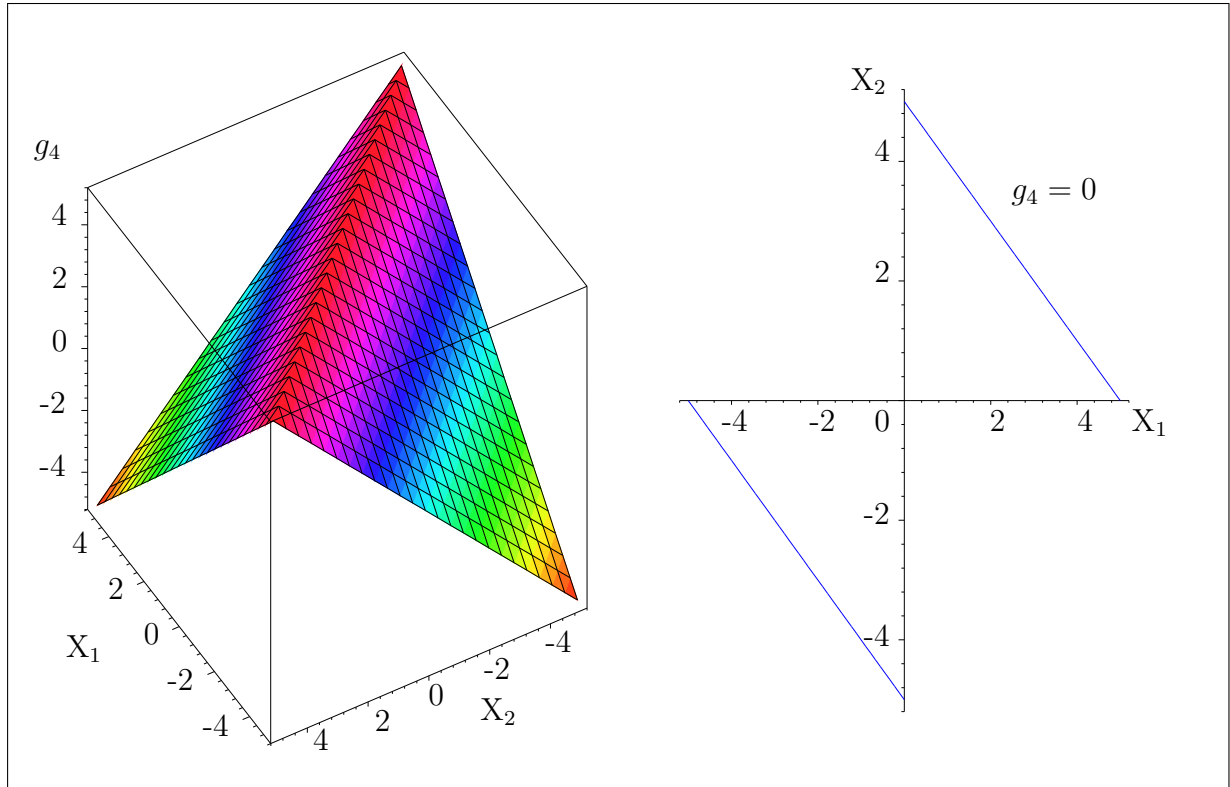
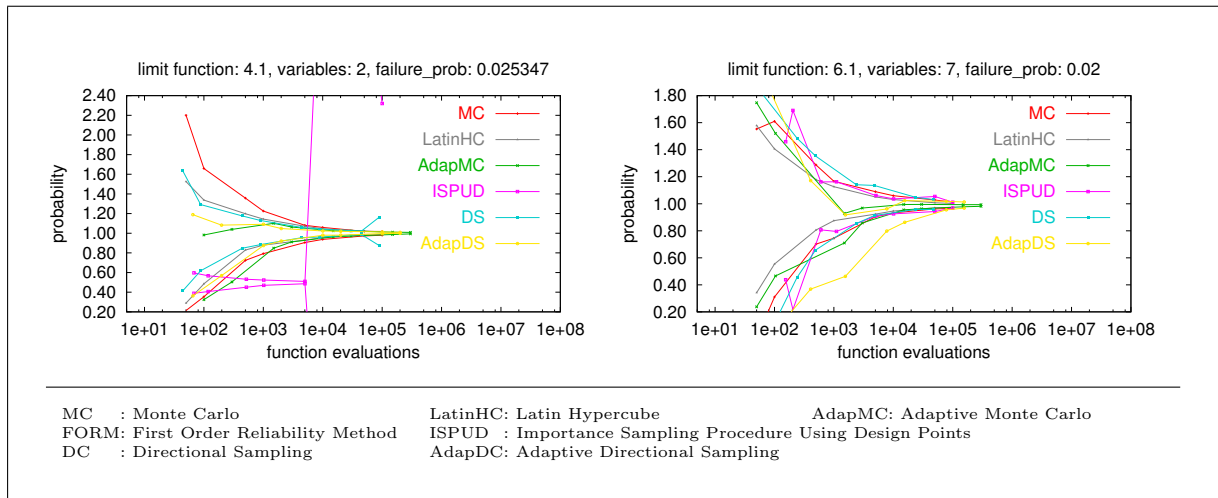
**Figure 14:** required number of samples for  $\text{tol}=0.5$  for varying dimension (number of random variables)

one hand due to the fact, that with the applied DOE scheme the design point is always present in the response surface, and on the other hand response surfaces are a good approximation of the real curve in low dimensions. For both response surface methods a strong dependence on the dimension is observed. For 5 random variables the accuracy of the estimated probability is within the range of the other methods. Further investigation with a higher number of random variables (with this example) shows, that in the case of a positive curvatures (LSF 2.2 and 2.4) the response surfaces can give accurate estimates of the exact probability up to 10 variables, whereas, as already observed in fig.13 with 5 variables, the application of response surfaces for negative curvatures should be limited to 5 dimensions.

Furthermore it is observed, that directional sampling gives precise estimates, especially in the case of positive curvatures. This is due to the fact, that the assumption for the failure surface for directional sampling (hyperspheres) correspond to this cases. This is further illustrated in fig. 14, where for a given tolerance  $\text{tol}=0.5$  the number of required samples is plotted as a function of the probability level, so that

$$1 - \text{tol} > \frac{\mu - \sigma}{\mu_{\text{exact}}} < \frac{\mu + \sigma}{\mu_{\text{exact}}} < 1 + \text{tol}. \quad (70)$$

For the limit state functions 2.1 and 2.3 directional sampling does not give accurate results for higher dimensions, whereas for the limit state functions 2.2 and 2.4 precise results even with a large number of 50 variables can be obtained. Importance Sampling Using (the) Design Point is in general the method of choice in this example. This is due to the fact, that the design point could be determined correctly within the optimization with NLPQL.


 Figure 15: limit state function  $g_4$ 

 Figure 16: interval of the obtained failure probability  $[\mu \pm \sigma]$  for limit functions 4 and 6

## Limit state function 4

$$g_4 = 5 - |X_1 + X_2| = 0 \quad (71)$$

$X_1$  : independent normal distributed  $\mu = 0, \sigma^2 = 1$

$X_2$  : independent normal distributed  $\mu = 0, \sigma^2 = 2$

The possibility to accurately capture multiple  $\beta$ -points is tested with a roof-like function. The results are illustrated in fig.16. Obviously the optimization with NLPQL determines only one design point and as a result the approximation with FORM gives only half of the exact failure probability. ISPUD samples with the design point as mean value. Until a certain number of samples no sample in the failure domain corresponding to the

opposite limit function is observed. That is the reason for ISPUD to converge towards 0.5. At about 5000 samples also samples coresponding to the opposite limit function are present, but this leads to a high variance of the estimator. The adaptive sampling procedure can deal with two design points, but for a low number of samples the correlation matrix in the first iteration step could not be determined accurately and as a result the sampling density function in the next iteration steps deviate from the optimal sampling density function, which leads in general to an underestimation of the failure probability. Directional Sampling is not influenced by the presence of multiple design points, and in the limit case of a hyper sphere as limit state function it always gives accurate results.

### 2.3 Limit state function 6

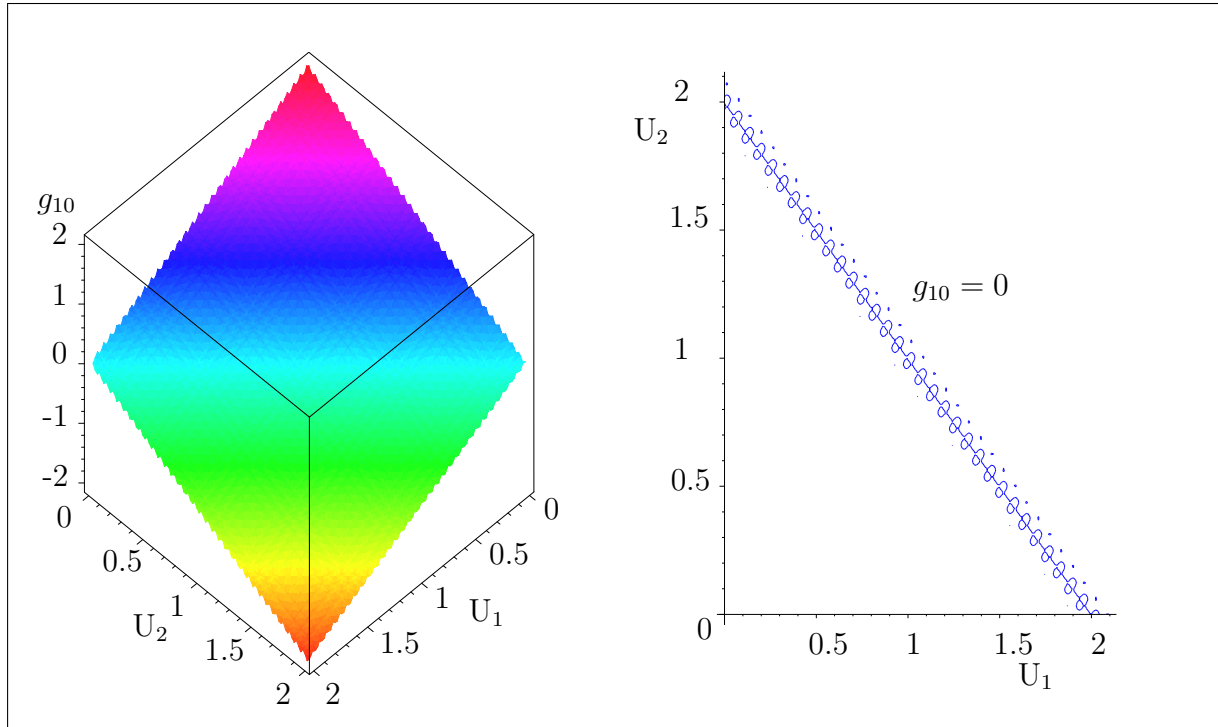
$$\begin{aligned}
g_6 &= \min(g_{6a}, g_{6b}, g_{6c}) = 0 \\
g_{6a} &= X_1 + 2X_3 + 2X_4 + X_5 - 5X_6 - 5X_7 \\
g_{6b} &= X_1 + 2X_2 + X_4 + X_5 - 5X_6 \\
g_{6c} &= X_2 + 2X_3 + X_4 - 5X_7 \\
X_{1..5} &: \text{independent lognormal distributed } \mu = 60, \sigma^2 = 6 \\
X_6 &: \text{independent Gumbel distributed } \mu = 20, \sigma^2 = 6 \\
X_7 &: \text{independent Gumbel distributed } \mu = 25, \sigma^2 = 7.5
\end{aligned} \tag{72}$$

The limit state function  $g_6$  is used to verify, if the methods can handle series systems. The failure surface corresponds to 3 hyperplanes. The results are illustrated in fig.16. With FORM a relative failure probability of 0.631 with 96 limit state function evaluations is obtained. This discrepancy to the correct solution is due to the fact, that FORM can only accurately deal with limitstate functions consisting of one hyperplane. ISPUD with only a few samples can already give accurate results, but as further illustrated in example 11, it is essential, that the mean value of the first sampling density ( $\beta$ -point) is close to the regions, that mainly contribute to the result. It is further obvious, that a certain number of samples is required to accurately perform an adaptation (Adaptive Sampling and Adaptive Directional Sampling).

### 2.4 Limit state function 10

$$\begin{aligned}
g_{10} &= 2 - (U_1 + U_2) + 0.05 [\sin(100U_1) + \cos(100U_2)] \\
U_i &: \text{independent normal distributed } \mu = 0, \sigma^2 = 1
\end{aligned} \tag{73}$$

A noisy limit state function is tested by adding to a hyperplane a trigonometric term. The factor 0.05 representing the noise is relatively high, but in order to test the influence on the methods a worst case scenario has been favored. An illustration of the failure surface is given in fig.17. Obviously the failure surface is not convex and bubbles within the safe region corresponding to regions of failure are present. As a result all the directional sampling methods (and furthermore the response surface methods) are theoretically not applicable, but in practice the distance between the points with  $g_{10} = 0$  for a given sampling direction is so close, that accurate results can be obtained. All the directional methods should theoretically (for a high number of sample directions) result in a failure probability that is higher than the correct result, since the first point with  $g_{10} = 0$  for a given direction is searched. Due to the implementation, where a line search algorithm



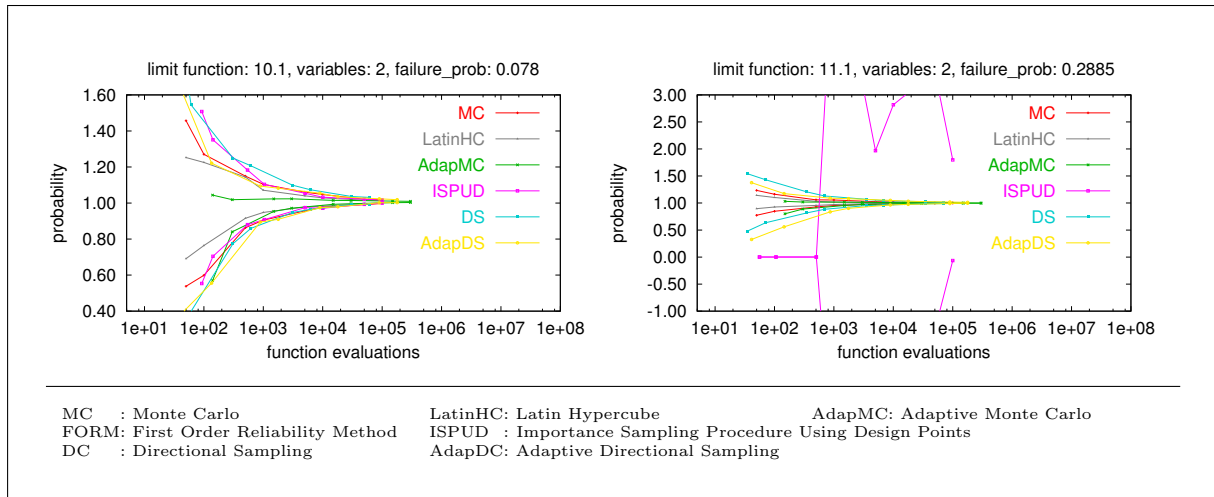
**Figure 17:** limit state function  $g_{10}$

is applied, it is not assured, that the first point is found, but any of the points with  $g_{10} = 0$  are possible. This averaging finally leads to the correct estimation of the failure probabilities, illustrated in fig.18. The optimization with NLPQL, which is a gradient based algorithm, sticks to a local minimum, which results in an estimate 5.21 of the relative failure probability for FORM. For lower noise factors the optimization did determine the exact  $\beta$ -point and an estimation with FORM yielded accurate results. Due to the high failure probability the number of samples for the first iteration step within the failure domain allowed an accurate estimation of the covariance matrix and as a result the Adaptive Sampling Procedure is for this limit state function superior to all other methods. It is further observed, that Latin Hypercube yields better estimates of the probability than plain Monte Carlo. The estimated  $\beta$ -point from NLPQL is sufficiently close to the exact one, so that the Importance Sampling Procedure Using (the) Design Point still gives accurate results.

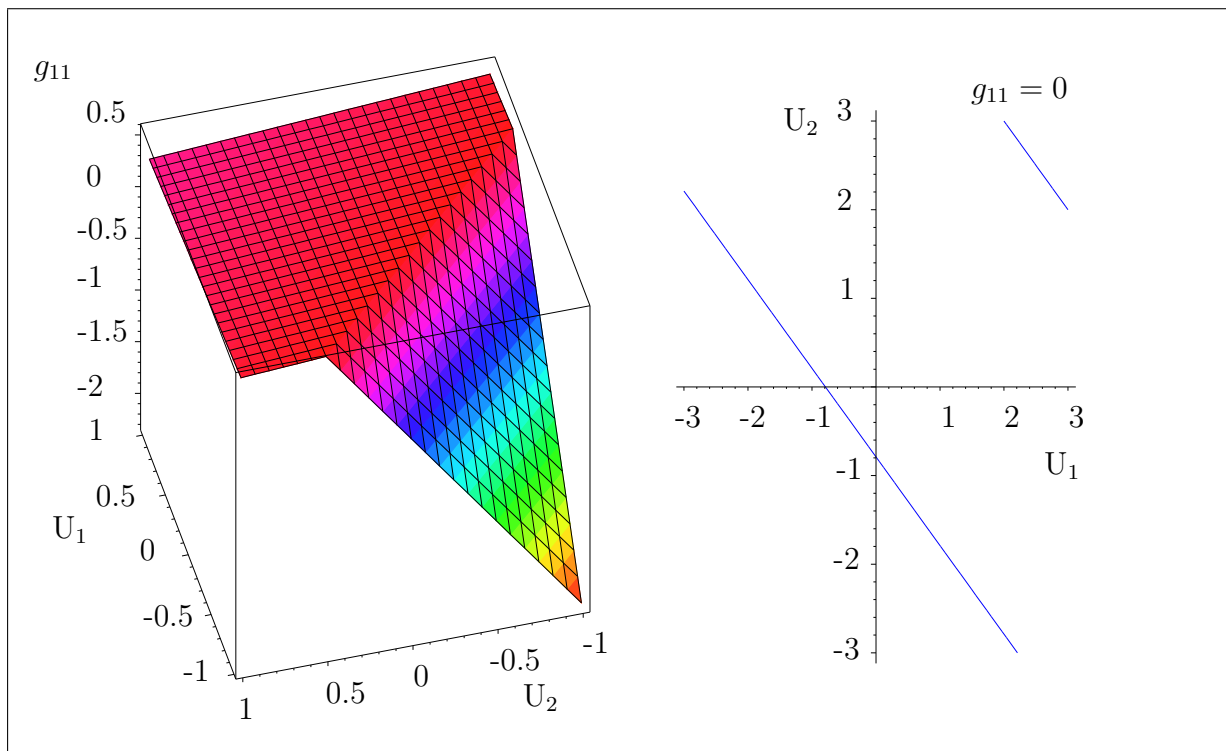
## 2.5 Limit state function 11

$$\begin{aligned} g_{11} &= 1 - |U_1 + U_2 + 0.5| + 0.9(U_1 + U_2) \\ U_i &: \text{independent normal distributed } \mu = 0, \sigma^2 = 1 \end{aligned} \quad (74)$$

The final example is equivalent to example 4, with the difference, that the two hyperplanes have different a different slope. This situation occurs, if a problem is nearly linear within a certain range, but at a certain point non-linearities lead to an abrupt modification of the response. Obviously a gradient based optimization must fail for this situation, which leads in this case to a strong underestimation of the failure probability with FORM (relative failure probability  $5 \cdot 10^{-5}$ ). Since the distance between the real  $\beta$ -point and the approximation is very large, ISPUD also fails for this example. The high failure probability leads to accurate estimates with Adaptive Sampling, Latin Hypercube and



**Figure 18:** intervall of the obtained failure probability  $[\mu \pm \sigma]$  for limit functions 10 and 11



**Figure 19:** limit state function  $g_{11}$

Monte Carlo, but accurate estimators are obtained for all methods.



### 3 Summary and Conclusion

In this paper different methods for the calculation of probabilities are investigated with respect to several parameters as the number of random variables, the probability level, the ability to deal with multiple design points, curvature and the shape of the limit state function and the existence of noise in the response function. As the different influencing parameters have a varying influence on the presented methods, there is for a specific problem generally a method, that is superior to all others. It is however difficult to determine this method exactly, since the limit state function and the probability level is not known before this procedure is applied. The choice of the appropriate method has to be based on estimates of the influencing parameters from engineering experience. If certain parameters are found to diverge from the assumed values, possibly a recalculation with a different method is necessary.

The most robust methods are Monte Carlo and Latin Hypercube, whereas the latter was in general slightly more accurate. The main drawback of these methods is, that the variance of the estimator of probabilities is quadratically dependent on the probability level, and as a result for low probabilities a large number of samples is required. An alternative is the use of importance sampling techniques. The first investigated method is the Adaptive Sampling approach, where after a first Monte Carlo simulation, possibly with adapted variance, the sampling density is determined from the samples of the previous iteration step within the failure domain. The method gave accurate results, if the number of samples within the failure domain was sufficient to accurately determine the correlation, but this number was relatively large, especially in higher dimensions, and as a result the method was for most examples not appropriate. A second importance Sampling technique is ISPUD (Importance Sampling Using Design Point). The sampling is centered near the design point, which is determined in a previous optimization step. This method was in many cases superior to all other methods, even if the design point could not be determined exactly. Only for certain problems with a sudden failure as in example 11 the estimate of the probability was rather erroneous. A rather low dependence on the dimension of the problem was observed. The method is a good choice for problems, where the design point can be accurately determined by an optimization (gradient based optimizer NLPQL).

A linearization of the limit state function at the design point with FORM (First Order Reliability Method) gives often accurate results and at least an idea about the expected failure probabilities and for linear problems the exact probability is obtained. FORM requires the least number of samples compared to all other methods, but it is not able to deal with multiple design points and the determination of the design point has to be possible. For certain problems with a noisy response the gradient based optimizer has failed, and for highly curved limit state function the linear assumption is not reasonable. Finally methods based on response surfaces have been tested. It has been observed, that for a small number of random variables ( $n < 5$ ) the response surface methods are accurate, even for small failure probabilities. The number of required samples is exponentially increasing with the number of samples, and in higher dimensions ( $n > 10$ ) the response surfaces are not applicable any more. An advantage of the methods based on response surfaces is, that the response surface itself can be investigated by the engineer.



As a conclusion the following proposal for the method of choice for a specific problem are made:

- design point determinable by a gradient based optimization
  - linear problems: FORM
  - nonlinear problems : ISPUD
- low number of random variables ( $< 5 - 10$ )
  - response surfaces
- high probability level ( $> 10^{-3}$ )
  - Latin Hypercube

## References

- ABRAMOWITSCH, M. ; STEGUN, I.M.: *Handbook of Mathematical Functions*. New York, USA : Dover Publications, 1972
- BJERAGER, P.: Probability Integration by Directional Simulation. In: *Journal of Engineering Mechanics, ASCE* Vol. 114 (1988), Nr. No. 8, S. 1285 – 1302
- BOURGUND, U. ; BUCHER, C. G.: Importance Sampling Procedure Using Design Points (ISPUD) - a User's Manual. In: *Bericht Nr. 8-86*. Innsbruck : Institut für Mechanik, Universität Innsbruck, 1986
- BREITUNG, K. W.: Probability Approximations by Log Likelihood Maximization. In: *Journal of Engineering Mechanics, ASCE* Vol. 117 (1991), Nr. No. 3, S. 457 – 477
- BUCHER, C. G.: *Adaptive Sampling - An Iterative Fast Monte Carlo Procedure*. Vol. 5, No. 2. Structural Safety, 1988
- DITLEVSEN, O. ; BJERAGER, P.: Plastic Reliability Analysis by Directional Simulation. In: *Journal of Engineering Mechanics, ASCE* Vol. 115 (1989), Nr. No. 6, S. 1347 – 1362
- HASOFER, A. M. ; LIND, N. C.: Exact and Invariant Second-Moment Code Format. In: *Journal of the Engineering Mechanics Division, ASCE* Vol. 100 (1974)
- HOHENBICHLER, M. ; RACKWITZ, R.: Improvement of Second-Order Reliability Estimates by Importance Sampling. In: *Journal of Engineering Mechanics, ASCE* Vol. 114, No. 12 (1988), S. 2195 – 2198
- IMAN, R. L. ; CONOVER, W. J.: *A Distribution Free Approach to Inducing Rank Correlation Among Input Variables*. Vol. B11. Communications in Statistics, 1982
- McKAY, M.D. ; BECKMAN, R.J. ; CONOVER, W.J.: A comparison of three methods for selecting values of input variables in the analysis of output from a computer code. In: *Technometrics* 21 (1979), Nr. 2, S. 239–245

## REFERENCES

- RUBINSTEIN, R. Y.: *Simulation and the Monte Carlo Method*. New York, USA : John Wiley & Sons, 1981
- SCHITTKOWSKI, Klaus: NLPQL: A Fortran subroutine for solving constrained nonlinear programming problems. In: *Annals of Operations Research* 5 (1985), S. 485–500
- SHINOZUKA, M.: Basic Analysis of Structural Safety. In: *Journal of Structural Engineering, ASCE* Vol. 109 (1983), Nr. No. 3, S. 721 – 740
- STEIN, M.L.: Large sample properties of simulations using latin hypercube sampling. In: *Technometrics* 29 (1987), Nr. 2, S. 143–151
- TVEDT, L.: *Two second-order approximations to the failure probability*. Oslo, Norway : Veritas Report R74-33, Det Norske Veritas, 1983