

Appendix A: The classical AK-MCS method

The Kriging method regards the g-function as an n -dimensional Gaussian random field, which reads:

$$g(\mathbf{x}) = \mathbf{f}(\mathbf{x})^T \boldsymbol{\beta} + Z(\mathbf{x}) \quad (\text{A1})$$

where $\mathbf{f}(\mathbf{x})^T \boldsymbol{\beta}$ refers to a linear regression model with a set of functional basis $\mathbf{f} = [f_1, \dots, f_d]$,

$Z(\mathbf{x})$ indicates a n -dimensional Gaussian random field with zero mean and variance σ_g^2 . The

covariance function of the random field for any two point $\mathbf{x}_i$ and $\mathbf{x}_j$ is given as:

$$C_{gg}(\mathbf{x}_i, \mathbf{x}_j) = \sigma_g^2 R(\mathbf{x}_i, \mathbf{x}_j) \quad (\text{A2})$$

where $R(\mathbf{x}_i, \mathbf{x}_j)$ is the autocorrelation function, and the most commonly used one is the

exponential one given by:

$$R(\mathbf{x}_i, \mathbf{x}_j) = \exp\left(-\sqrt{\sum_{k=1}^n \frac{(x_{ki} - x_{kj})^2}{\lambda_k^2}}\right) \quad (\text{A3})$$

and λ_k is a parameter measuring the strength of the autocorrelation. Then given a set of N_0 training samples

$\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(N_0)}$ and the corresponding values $z^{(1)}, \dots, z^{(N_0)}$ of the

gfunction, the mean value $\mu_g(\mathbf{x})$ and the variance $\sigma_g^2(\mathbf{x})$ for a new point $\mathbf{x}$ conditional on these training samples are given as [24]:

$$\begin{aligned} \mu_g(\mathbf{x}) &= \mathbf{f}(\mathbf{x})^T \boldsymbol{\beta} + \mathbf{r}(\mathbf{x})^T \mathbf{R}^{-1}(\mathbf{z} - \mathbf{F}\boldsymbol{\beta}) \\ \sigma_g^2(\mathbf{x}) &= \sigma_g^2 \left(1 - \mathbf{r}(\mathbf{x})^T \mathbf{R}^{-1} \mathbf{r}(\mathbf{x}) - \mathbf{u}(\mathbf{x})^T \mathbf{F} \mathbf{R} \mathbf{F}^{-1} \mathbf{u}(\mathbf{x}) \right) \end{aligned} \quad (\text{A4})$$

where $\mathbf{R}$ is the correlation matrix of these N_0 training samples, $\mathbf{r}(\mathbf{x})$ refers to the vector of correlation functions between $\mathbf{x}$ and the N_0 training samples, $\mathbf{F}$ indicates the regression matrix

computed by $F_{ij} = f_j(\mathbf{x}^{(i)})$ with $i = 1, \dots, N$ and $j = 1, \dots, d$, $\mathbf{u}(\mathbf{x}) = \mathbf{F}^T \mathbf{R}^{-1} \mathbf{r}(\mathbf{x}) - \mathbf{f}(\mathbf{x})$ and

the least square estimates of the regression coefficients are computed by $\hat{\boldsymbol{\beta}} = (\mathbf{F}^T \mathbf{R}^{-1} \mathbf{F})^{-1} \mathbf{F}^T \mathbf{R}^{-1} \mathbf{z}$.

Then $\square_g \square \mathbf{x}$ refers to the Kriging surrogate model, and $\square_{g^2} \square \mathbf{x}$ indicates the mean square error

for $j = 1, 2, \dots, N_0$

Step 1: Generate a MC population $\mathbf{W}$ of N (N is large, i.e., $1e5$) samples of the input vector. Randomly chose N_0 (N_0 is small, i.e., 12) samples from the population $\mathbf{W}$. Estimate the corresponding g function values for these N_0 training samples, and attribute these N_0 samples to the training population $\mathbf{W}_t$.

Step 3: Generate the Kriging prediction $\hat{\mathbf{x}}_{g^1}^{\mathbf{x}_{\mathbf{f}_j}}$ and $\hat{\mathbf{x}}_{g^2}^{\mathbf{x}_{\mathbf{f}_j}}$ for each of the remaining N -

□. If $\min_{j \in 1}^N U_j \leq 2$, then add the sample with the minimum U value to the training population

Appendix B: IS Estimators of the GRS indices

Let x_j^i indicate the random replication of x_j ($j = 1^2, \dots, n$), $\mathbf{x}_i^i = [x_{i1}^i, \dots, x_{in}^i]$, $\mathbf{X} = [X_1, \dots, X_n]$, $\mathbf{X}_i = [x_{i1}^i, \dots, x_{in}^i]$,

$$\begin{array}{c} \int_{R_n} I_F \, dx \, p_{X_k} \, dx_k \\ \int_{R_n} I_F \, dx' \, p_{X_k} \, dx_{k'} \end{array}$$

$$\int_{n_1}^{n_2} h(x) dx = \int_{n_1}^{n_2} h(x) dx$$

$$h(x_k) - \int_{x_k}^{x_{k'}} \frac{1}{x} dx = \int_{x_k}^{x_{k'}} \frac{1}{x} dx - \int_{x_k}^{x_{k'}} \frac{1}{x} dx$$

$$\begin{aligned} & \left\| \left(\sum_{k \in \mathbb{Z}^n} \left| \hat{f}_k \right|^2 \right)^{1/2} \right\|_{L^p(\mathbb{R}^n)} \leq C \left\| \left(\sum_{k \in \mathbb{Z}^n} \left| \hat{f}_k \right|^2 \right)^{1/2} \right\|_{L^p(\mathbb{R}^n)} \\ & \left\| \left(\sum_{k \in \mathbb{Z}^n} \left| \hat{f}_k \right|^2 \right)^{1/2} \right\|_{L^p(\mathbb{R}^n)} \leq C \left\| \left(\sum_{k \in \mathbb{Z}^n} \left| \hat{f}_k \right|^2 \right)^{1/2} \right\|_{L^p(\mathbb{R}^n)} \end{aligned}$$

$$\begin{array}{c} \square \\ \square \square \quad k \square \quad k \square \square \square \quad k \square \quad k \square \quad k \square \square \end{array}$$

$$\begin{array}{c} n \quad \text{---} \quad p \, x_k \square \quad k \square \square \quad n \quad \text{---} \quad p \, x_k \square \quad k' \square \quad n p \, x_k \square \quad k' \square \quad n \square \, x \, h \end{array}$$

$$\begin{array}{c} F \quad \square \quad \square I \, x \quad \square \\ \square \square_{R2n} I_F \square \square \, x \square \quad k \square 1 \, h \, x_k \square \quad k \square \square \square I \square \, x \square' i \square \quad k \square \square 1, k \, i \, h \, x_k \square \quad k' \square \square_F \square \quad , \square \square \square k \square 1 \, h \, x_k \square \quad k' \square \square \square \square \square k \square 1 \square \quad k \square \quad k \square \square \, x_k \square \end{array}$$

$$\square d \, x \, x_k d \, k' \square \square$$

$$\begin{array}{c} \square \quad n p \, x_k \square \quad k \quad \text{---} \quad \text{---} \quad \square \square \quad , \quad n \quad p \, x_k \square \quad k' \square \end{array}$$

$$\begin{array}{c} n p \, x_k \square \quad k' \square \square \square \square \\ \square \\ \square \, E_h \square I_F \square \square \, x \square \quad \text{---} \quad \square I \end{array}$$

$$\square \square \square k \square 1 \, h \, x_k \square \quad k \square \square \square_F \square \, x \square i \square \quad k \square \square \square 1, k \, i \, h \, x_k \square \quad k' \square \square \square I_F \square \, x \square \square \square k \square 1 \, h \, x_k \square \quad k' \square \square \square \square \square \square$$

(A5)

Then the sample matrices $\mathbf{B}$, $\mathbf{A}$ and $\mathbf{C}$ introduced in subsection 3.4 can be regarded as the IS samples of the random vectors $\mathbf{x}$, $\mathbf{x}'$ and $\mathbf{x}_{\square' i}$, and the IS estimator of V_i is given as:

V_i

IS estimator of V_{Ti} is derived as:

$\mathbf{x}_i'$, and the

$$V_{Ti} = \frac{1}{N} \sum_{i=1}^N \frac{1}{\sigma_i^2} \left[\sum_{j=1}^n \frac{1}{\sigma_j^2} \left(\frac{\partial g}{\partial x_j} \right)^2 + \sum_{j=1}^n \frac{1}{\sigma_j^2} \left(\frac{\partial g}{\partial x_j} \right) \left(\frac{\partial g}{\partial x_i} \right) \right] \quad (A8)$$

Appendix C: FORM estimators of the GRS indices

Assume that the random input vector follows independent Gaussian distribution with mean vector $\boldsymbol{\mu} = [\mu_1, \mu_2, \dots, \mu_n]^T$ and SD vector $\boldsymbol{\sigma} = [\sigma_1, \sigma_2, \dots, \sigma_n]^T$. The first order Taylors series of the limit state

function expanded at the MPP $\mathbf{x}^* = [x_1^*, x_2^*, \dots, x_n^*]$ is given as follows:

$$z = \mathbf{x}^T \mathbf{g} \mathbf{x}^* + \sum_{j=1}^n a_j (x_j - x_j^*) \quad (A9)$$

where a_j indicates the partial derivatives of the g-function w.r.t. x_j at the MPP, and

$$a_0 = g(\mathbf{x}^*)$$

Then, based on the first line of Eq. (A5), the main partial variance V_i can be derived as:

$$V_i = \sum_{j=1}^n \frac{1}{\sigma_j^2} \left(\frac{\partial g}{\partial x_j} \right)^2 + \sum_{j=1}^n \frac{1}{\sigma_j^2} \left(\frac{\partial g}{\partial x_j} \right) \left(\frac{\partial g}{\partial x_i} \right) \quad (A10)$$

$$\Pr[z = 0] = \Pr[\mathbf{x} = \mathbf{x}^*] = 0$$

where $z = \mathbf{x}^T \mathbf{g} \mathbf{x}^*$ and $z = \mathbf{x}^T \mathbf{g} \mathbf{x}^*$ are two linear function of the vector $\mathbf{x}$ and $\mathbf{x}^*$, respectively, thus can be

regarded as two correlated Gaussian random variables with covariance $\sigma_{zm}^2 = \sigma_i^2 \sigma_m^2$. $z = \mathbf{x}^T \mathbf{g} \mathbf{x}^*$ and

$z = \mathbf{x}^T \mathbf{g} \mathbf{x}^*$ have the same mean value $\mu_z = a_0$ and the same SD $\sigma_z = \sqrt{\sum_{j=1}^n a_j^2 \sigma_j^2}$. Then

V_i can be estimated by:

$$\hat{V}_i = \Phi_2[0,0; \sigma_m, \sigma_m] \hat{P}_f^2 \quad (A11)$$

where $\Phi_2[0,0; \sigma_m, \sigma_m]$ indicates the bivariate joint CDF of Gaussian distribution with mean vector

$$\sigma_m^2 = \sigma_{z2m2}^2$$

σ_m^2 and covariance matrix σ_m^2 calculated at the point [0,0].

$$\sigma_m^2 = \sigma_{zm}^2$$

Similarly, based on the first line of Eq.(A7), the total partial variance V_{Ti} can be approximated by:

$$V_{Ti} = P_f \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \left(\sum_{j=1}^n \frac{\partial f}{\partial x_j} \right)^2 p_j(x_j) dx_j \quad (A12)$$

$$= P_f \Pr \{ z = 0 \mid z = x_i' \}$$

where $z = x$ and $z = x_i'$ are linear functions of x and x_i' , thus are also two correlated Gaussian

random variables with covariance $\sigma_{z^2}^2 = \sigma_{j=1, j}^2 a_j^2$. Then, based on Eq. (A12), the total partial

variance V_{Ti} can be further derived as:

$$\hat{V}_{Ti} = \hat{P}_f \sigma_{z^2}^2 [0,0]; \sigma_t, \sigma_t \quad (A13)$$

where $\sigma_t = \sigma_m$ and $\sigma_t = \sigma_{z22} \sigma_{z22}$.

