

SANDIA REPORT

SAND2014-4253

Unlimited Release

Printed May 2014; updated May 21, 2014

Dakota, A Multilevel Parallel Object-Oriented Framework for Design Optimization, Parameter Estimation, Uncertainty Quantification, and Sensitivity Analysis: Version 6.0 Theory Manual

Brian M. Adams, Lara E. Bauman, William J. Bohnhoff, Keith R. Dalbey,
John P. Eddy, Mohamed S. Ebeida, Michael S. Eldred, Patricia D. Hough,
Kenneth T. Hu, John D. Jakeman, Laura P. Swiler, J. Adam Stephens,
Dena M. Vigil, Timothy M. Wildey

Prepared by
Sandia National Laboratories
Albuquerque, New Mexico 87185 and Livermore, California 94550

Sandia National Laboratories is a multi-program laboratory managed and operated by Sandia Corporation, a wholly owned subsidiary of Lockheed Martin Corporation, for the U.S. Department of Energy's National Nuclear Security Administration under contract DE-AC04-94AL85000.

Approved for public release; further dissemination unlimited.



Issued by Sandia National Laboratories, operated for the United States Department of Energy by Sandia Corporation.

NOTICE: This report was prepared as an account of work sponsored by an agency of the United States Government. Neither the United States Government, nor any agency thereof, nor any of their employees, nor any of their contractors, subcontractors, or their employees, make any warranty, express or implied, or assume any legal liability or responsibility for the accuracy, completeness, or usefulness of any information, apparatus, product, or process disclosed, or represent that its use would not infringe privately owned rights. Reference herein to any specific commercial product, process, or service by trade name, trademark, manufacturer, or otherwise, does not necessarily constitute or imply its endorsement, recommendation, or favoring by the United States Government, any agency thereof, or any of their contractors or subcontractors. The views and opinions expressed herein do not necessarily state or reflect those of the United States Government, any agency thereof, or any of their contractors.

Printed in the United States of America. This report has been reproduced directly from the best available copy.

Available to DOE and DOE contractors from
U.S. Department of Energy
Office of Scientific and Technical Information
P.O. Box 62
Oak Ridge, TN 37831

Telephone: (865) 576-8401
Facsimile: (865) 576-5728
E-Mail: reports@adonis.osti.gov
Online ordering: <http://www.osti.gov/bridge>

Available to the public from
U.S. Department of Commerce
National Technical Information Service
5285 Port Royal Rd
Springfield, VA 22161

Telephone: (800) 553-6847
Facsimile: (703) 605-6900
E-Mail: orders@ntis.fedworld.gov
Online ordering: <http://www.ntis.gov/help/ordermethods.asp?loc=7-4-0#online>



Dakota, A Multilevel Parallel Object-Oriented Framework for
Design Optimization, Parameter Estimation, Uncertainty
Quantification, and Sensitivity Analysis:
Version 6.0 Theory Manual

Brian M. Adams, Mohamed S. Ebeida, Michael S. Eldred, John D. Jakeman,
Laura P. Swiler, J. Adam Stephens, Dena M. Vigil, Timothy M. Wildey
Optimization and Uncertainty Quantification Department

William J. Bohnhoff
Radiation Transport Department

John P. Eddy
System Readiness and Sustainment Technologies Department

Kenneth T. Hu
Validation and Uncertainty Quantification Department

Keith R. Dalbey
Mission Analysis and Simulation Department

Lara E. Bauman, Patricia D. Hough
Quantitative Modeling and Analysis Department

Sandia National Laboratories
P.O. Box 5800
Albuquerque, NM 87185

Abstract

The Dakota (Design Analysis Kit for Optimization and Terascale Applications) toolkit provides a flexible and extensible interface between simulation codes and iterative analysis methods. Dakota contains algorithms for optimization with gradient and nongradient-based methods; uncertainty quantification with sampling, reliability, and stochastic expansion methods; parameter estimation with nonlinear least squares methods; and sensitivity/variance analysis with design of experiments

and parameter study methods. These capabilities may be used on their own or as components within advanced strategies such as surrogate-based optimization, mixed integer nonlinear programming, or optimization under uncertainty. By employing object-oriented design to implement abstractions of the key components required for iterative systems analyses, the Dakota toolkit provides a flexible and extensible problem-solving environment for design and performance analysis of computational models on high performance computers.

This report serves as a theoretical manual for selected algorithms implemented within the Dakota software. It is not intended as a comprehensive theoretical treatment, since a number of existing texts cover general optimization theory, statistical analysis, and other introductory topics. Rather, this manual is intended to summarize a set of Dakota-related research publications in the areas of surrogate-based optimization, uncertainty quantification, and optimization under uncertainty that provide the foundation for many of Dakota's iterative analysis capabilities.

Contents

1	Reliability Methods	9
1.1	Local Reliability Methods	9
1.1.1	Mean Value	9
1.1.2	MPP Search Methods	10
1.1.2.1	Limit state approximations	12
1.1.2.2	Probability integrations	14
1.1.2.3	Hessian approximations	14
1.1.2.4	Optimization algorithms	15
1.1.2.5	Warm Starting of MPP Searches	15
1.2	Global Reliability Methods	15
1.2.1	Importance Sampling	16
1.2.2	Efficient Global Optimization	16
1.2.2.1	Gaussian Process Model	17
1.2.2.2	Expected Improvement Function	18
1.2.2.3	Expected Feasibility Function	19
2	Stochastic Expansion Methods	21
2.1	Orthogonal polynomials	21
2.1.1	Askey scheme	21
2.1.2	Numerically generated orthogonal polynomials	22
2.2	Interpolation polynomials	22
2.2.1	Nodal interpolation	22
2.2.1.1	Global value-based	23
2.2.1.2	Global gradient-enhanced	23
2.2.1.3	Local value-based	24

2.2.1.4	Local gradient-enhanced	24
2.2.2	Hierarchical interpolation	24
2.3	Generalized Polynomial Chaos	25
2.3.1	Expansion truncation and tailoring	26
2.4	Stochastic Collocation	27
2.4.1	Value-Based Nodal	27
2.4.2	Gradient-Enhanced Nodal	28
2.4.3	Hierarchical	28
2.5	Transformations to uncorrelated standard variables	29
2.6	Spectral projection	30
2.6.1	Sampling	31
2.6.2	Tensor product quadrature	31
2.6.3	Smolyak sparse grids	32
2.6.4	Cubature	33
2.7	Linear regression	34
2.7.1	Cross validation	36
2.7.2	Orthogonal Least Interpolation	36
2.8	Analytic moments	37
2.9	Local sensitivity analysis: derivatives with respect to expansion variables	38
2.10	Global sensitivity analysis: variance-based decomposition	38
2.11	Automated Refinement	39
2.11.1	Uniform refinement with unbiased grids	40
2.11.2	Dimension-adaptive refinement with biased grids	40
2.11.3	Goal-oriented dimension-adaptive refinement with greedy adaptation	41
2.12	Multifidelity methods	42
3	Epistemic Methods	45
3.1	Dempster-Shafer theory of evidence (DSTE)	45
4	Surrogate Models	47
4.1	Kriging and Gaussian Process Models	47
4.1.1	Kriging & Gaussian Processes: Function Values Only	47
4.1.2	Gradient Enhanced Kriging	52
5	Surrogate-Based Local Minimization	57

5.1	Iterate acceptance logic	59
5.2	Merit functions	60
5.3	Convergence assessment	61
5.4	Constraint relaxation	61
6	Optimization Under Uncertainty (OUU)	65
6.1	Reliability-Based Design Optimization (RBDO)	65
6.1.1	Bi-level RBDO	65
6.1.2	Sequential/Surrogate-based RBDO	67
6.2	Stochastic Expansion-Based Design Optimization (SEBDO)	67
6.2.1	Stochastic Sensitivity Analysis	67
6.2.1.1	Local sensitivity analysis: first-order probabilistic expansions	68
6.2.1.2	Local sensitivity analysis: zeroth-order combined expansions	69
6.2.1.3	Inputs and outputs	69
6.2.2	Optimization Formulations	70
6.2.2.1	Bi-level SEBDO	70
6.2.2.2	Sequential/Surrogate-Based SEBDO	71
6.2.2.3	Multifidelity SEBDO	71

Chapter 1

Reliability Methods

1.1 Local Reliability Methods

Local reliability methods include the Mean Value method and the family of most probable point (MPP) search methods. Each of these methods is gradient-based, employing local approximations and/or local optimization methods.

1.1.1 Mean Value

The Mean Value method (MV, also known as MVFOSM in [53]) is the simplest, least-expensive reliability method because it estimates the response means, response standard deviations, and all CDF/CCDF response-probability-reliability levels from a single evaluation of response functions and their gradients at the uncertain variable means. This approximation can have acceptable accuracy when the response functions are nearly linear and their distributions are approximately Gaussian, but can have poor accuracy in other situations. The expressions for approximate response mean μ_g , approximate response variance σ_g^2 , response target to approximate probability/reliability level mapping ($\bar{z} \rightarrow p, \beta$), and probability/reliability target to approximate response level mapping ($\bar{p}, \bar{\beta} \rightarrow z$) are

$$\mu_g = g(\mu_{\mathbf{x}}) \quad (1.1)$$

$$\sigma_g^2 = \sum_i \sum_j Cov(i, j) \frac{dg}{dx_i}(\mu_{\mathbf{x}}) \frac{dg}{dx_j}(\mu_{\mathbf{x}}) \quad (1.2)$$

$$\bar{z} \rightarrow \beta : \quad \beta_{\text{CDF}} = \frac{\mu_g - \bar{z}}{\sigma_g}, \quad \beta_{\text{CCDF}} = \frac{\bar{z} - \mu_g}{\sigma_g} \quad (1.3)$$

$$\bar{\beta} \rightarrow z : \quad z = \mu_g - \sigma_g \bar{\beta}_{\text{CDF}}, \quad z = \mu_g + \sigma_g \bar{\beta}_{\text{CCDF}} \quad (1.4)$$

respectively, where $\mathbf{x}$ are the uncertain values in the space of the original uncertain variables (“x-space”), $g(\mathbf{x})$ is the limit state function (the response function for which probability-response level pairs are needed), and β_{CDF} and β_{CCDF} are the CDF and CCDF reliability indices, respectively.

With the introduction of second-order limit state information, MVSOSM calculates a second-order mean as

$$\mu_g = g(\mu_{\mathbf{x}}) + \frac{1}{2} \sum_i \sum_j Cov(i, j) \frac{d^2 g}{dx_i dx_j}(\mu_{\mathbf{x}}) \quad (1.5)$$

This is commonly combined with a first-order variance (Equation 1.2), since second-order variance involves higher order distribution moments (skewness, kurtosis) [53] which are often unavailable.

The first-order CDF probability $p(g \leq z)$, first-order CCDF probability $p(g > z)$, β_{CDF} , and β_{CCDF} are related to one another through

$$p(g \leq z) = \Phi(-\beta_{\text{CDF}}) \quad (1.6)$$

$$p(g > z) = \Phi(-\beta_{\text{CCDF}}) \quad (1.7)$$

$$\beta_{\text{CDF}} = -\Phi^{-1}(p(g \leq z)) \quad (1.8)$$

$$\beta_{\text{CCDF}} = -\Phi^{-1}(p(g > z)) \quad (1.9)$$

$$\beta_{\text{CDF}} = -\beta_{\text{CCDF}} \quad (1.10)$$

$$p(g \leq z) = 1 - p(g > z) \quad (1.11)$$

where $\Phi()$ is the standard normal cumulative distribution function. A common convention in the literature is to define g in such a way that the CDF probability for a response level z of zero (i.e., $p(g \leq 0)$) is the response metric of interest. Dakota is not restricted to this convention and is designed to support CDF or CCDF mappings for general response, probability, and reliability level sequences.

With the Mean Value method, it is possible to obtain importance factors indicating the relative importance of input variables. The importance factors can be viewed as an extension of linear sensitivity analysis combining deterministic gradient information with input uncertainty information, *i.e.* input variable standard deviations. The accuracy of the importance factors is contingent of the validity of the linear approximation used to approximate the true response functions. The importance factors are determined as:

$$ImpFactor_i = \left(\frac{\sigma_{x_i}}{\sigma_g} \frac{dg}{dx_i}(\mu_{\mathbf{x}}) \right)^2 \quad (1.12)$$

1.1.2 MPP Search Methods

All other local reliability methods solve an equality-constrained nonlinear optimization problem to compute a most probable point (MPP) and then integrate about this point to compute probabilities. The MPP search is performed in uncorrelated standard normal space (“u-space”) since it simplifies the probability integration: the distance of the MPP from the origin has the meaning of the number of input standard deviations separating the mean response from a particular response threshold. The transformation from correlated non-normal distributions (x-space) to uncorrelated standard normal distributions (u-space) is denoted as $\mathbf{u} = T(\mathbf{x})$ with the reverse transformation denoted as $\mathbf{x} = T^{-1}(\mathbf{u})$. These transformations are nonlinear in general, and possible approaches include the Rosenblatt [85], Nataf [26], and Box-Cox [12] transformations. The nonlinear transformations may also be linearized, and common approaches for this include the Rackwitz-Fiessler [80] two-parameter equivalent normal and the Chen-Lind [19] and Wu-Wirsching [101] three-parameter equivalent normals. Dakota employs the Nataf nonlinear transformation which is suitable for the common case when marginal distributions and a correlation matrix are provided, but full joint distributions are not known¹. This transformation occurs in the following two steps. To transform between the original correlated x-space variables and correlated standard normals (“z-space”), a CDF matching condition is applied for each of the marginal distributions:

$$\Phi(z_i) = F(x_i) \quad (1.13)$$

where $F()$ is the cumulative distribution function of the original probability distribution. Then, to transform between correlated z-space variables and uncorrelated u-space variables, the Cholesky factor $\mathbf{L}$ of a modified

¹If joint distributions are known, then the Rosenblatt transformation is preferred.

correlation matrix is used:

$$\mathbf{z} = \mathbf{L}\mathbf{u} \quad (1.14)$$

where the original correlation matrix for non-normals in x-space has been modified to represent the corresponding “warped” correlation in z-space [26].

The forward reliability analysis algorithm of computing CDF/CCDF probability/reliability levels for specified response levels is called the reliability index approach (RIA), and the inverse reliability analysis algorithm of computing response levels for specified CDF/CCDF probability/reliability levels is called the performance measure approach (PMA) [93]. The differences between the RIA and PMA formulations appear in the objective function and equality constraint formulations used in the MPP searches. For RIA, the MPP search for achieving the specified response level $\bar{z}$ is formulated as computing the minimum distance in u-space from the origin to the $\bar{z}$ contour of the limit state response function:

$$\begin{aligned} &\text{minimize} && \mathbf{u}^T \mathbf{u} \\ &\text{subject to} && G(\mathbf{u}) = \bar{z} \end{aligned} \quad (1.15)$$

where $\mathbf{u}$ is a vector centered at the origin in u-space and $g(\mathbf{x}) \equiv G(\mathbf{u})$ by definition. For PMA, the MPP search for achieving the specified reliability level $\bar{\beta}$ or first-order probability level $\bar{p}$ is formulated as computing the minimum/maximum response function value corresponding to a prescribed distance from the origin in u-space:

$$\begin{aligned} &\text{minimize} && \pm G(\mathbf{u}) \\ &\text{subject to} && \mathbf{u}^T \mathbf{u} = \bar{\beta}^2 \end{aligned} \quad (1.16)$$

where $\bar{\beta}$ is computed from $\bar{p}$ using Eq. 1.8 or 1.9 in the latter case of a prescribed first-order probability level. For a specified generalized reliability level $\bar{\beta}^*$ or second-order probability level $\bar{p}$, the equality constraint is reformulated in terms of the generalized reliability index:

$$\begin{aligned} &\text{minimize} && \pm G(\mathbf{u}) \\ &\text{subject to} && \beta^*(\mathbf{u}) = \bar{\beta}^* \end{aligned} \quad (1.17)$$

where $\bar{\beta}^*$ is computed from $\bar{p}$ using Eq. 6.7 (or its CCDF complement) in the latter case of a prescribed second-order probability level.

In the RIA case, the optimal MPP solution $\mathbf{u}^*$ defines the reliability index from $\beta = \pm \|\mathbf{u}^*\|_2$, which in turn defines the CDF/CCDF probabilities (using Equations 1.6-1.7 in the case of first-order integration). The sign of β is defined by

$$G(\mathbf{u}^*) > G(\mathbf{0}) : \beta_{\text{CDF}} < 0, \beta_{\text{CCDF}} > 0 \quad (1.18)$$

$$G(\mathbf{u}^*) < G(\mathbf{0}) : \beta_{\text{CDF}} > 0, \beta_{\text{CCDF}} < 0 \quad (1.19)$$

where $G(\mathbf{0})$ is the median limit state response computed at the origin in u-space² (where $\beta_{\text{CDF}} = \beta_{\text{CCDF}} = 0$ and first-order $p(g \leq z) = p(g > z) = 0.5$). In the PMA case, the sign applied to $G(\mathbf{u})$ (equivalent to minimizing or maximizing $G(\mathbf{u})$) is similarly defined by either $\bar{\beta}$ (for a specified reliability or first-order probability level) or from a $\bar{\beta}$ estimate³ computed from $\bar{\beta}^*$ (for a specified generalized reliability or second-order probability level)

$$\bar{\beta}_{\text{CDF}} < 0, \bar{\beta}_{\text{CCDF}} > 0 : \text{maximize } G(\mathbf{u}) \quad (1.20)$$

$$\bar{\beta}_{\text{CDF}} > 0, \bar{\beta}_{\text{CCDF}} < 0 : \text{minimize } G(\mathbf{u}) \quad (1.21)$$

where the limit state at the MPP ($G(\mathbf{u}^*)$) defines the desired response level result.

²It is not necessary to explicitly compute the median response since the sign of the inner product $\langle \mathbf{u}^*, \nabla_{\mathbf{u}} G \rangle$ can be used to determine the orientation of the optimal response with respect to the median response.

³computed by inverting the second-order probability relationships described in Section 1.1.2.2 at the current $\mathbf{u}^*$ iterate.

1.1.2.1 Limit state approximations

There are a variety of algorithmic variations that are available for use within RIA/PMA reliability analyses. First, one may select among several different limit state approximations that can be used to reduce computational expense during the MPP searches. Local, multipoint, and global approximations of the limit state are possible. [33] investigated local first-order limit state approximations, and [34] investigated local second-order and multipoint approximations. These techniques include:

1. a single Taylor series per response/reliability/probability level in $\mathbf{x}$ -space centered at the uncertain variable means. The first-order approach is commonly known as the Advanced Mean Value (AMV) method:

$$g(\mathbf{x}) \cong g(\mu_{\mathbf{x}}) + \nabla_{\mathbf{x}} g(\mu_{\mathbf{x}})^T (\mathbf{x} - \mu_{\mathbf{x}}) \quad (1.22)$$

and the second-order approach has been named AMV²:

$$g(\mathbf{x}) \cong g(\mu_{\mathbf{x}}) + \nabla_{\mathbf{x}} g(\mu_{\mathbf{x}})^T (\mathbf{x} - \mu_{\mathbf{x}}) + \frac{1}{2} (\mathbf{x} - \mu_{\mathbf{x}})^T \nabla_{\mathbf{x}}^2 g(\mu_{\mathbf{x}}) (\mathbf{x} - \mu_{\mathbf{x}}) \quad (1.23)$$

2. same as AMV/AMV², except that the Taylor series is expanded in $\mathbf{u}$ -space. The first-order option has been termed the $\mathbf{u}$ -space AMV method:

$$G(\mathbf{u}) \cong G(\mu_{\mathbf{u}}) + \nabla_{\mathbf{u}} G(\mu_{\mathbf{u}})^T (\mathbf{u} - \mu_{\mathbf{u}}) \quad (1.24)$$

where $\mu_{\mathbf{u}} = T(\mu_{\mathbf{x}})$ and is nonzero in general, and the second-order option has been named the $\mathbf{u}$ -space AMV² method:

$$G(\mathbf{u}) \cong G(\mu_{\mathbf{u}}) + \nabla_{\mathbf{u}} G(\mu_{\mathbf{u}})^T (\mathbf{u} - \mu_{\mathbf{u}}) + \frac{1}{2} (\mathbf{u} - \mu_{\mathbf{u}})^T \nabla_{\mathbf{u}}^2 G(\mu_{\mathbf{u}}) (\mathbf{u} - \mu_{\mathbf{u}}) \quad (1.25)$$

3. an initial Taylor series approximation in $\mathbf{x}$ -space at the uncertain variable means, with iterative expansion updates at each MPP estimate ($\mathbf{x}^*$) until the MPP converges. The first-order option is commonly known as AMV+:

$$g(\mathbf{x}) \cong g(\mathbf{x}^*) + \nabla_{\mathbf{x}} g(\mathbf{x}^*)^T (\mathbf{x} - \mathbf{x}^*) \quad (1.26)$$

and the second-order option has been named AMV²+

$$g(\mathbf{x}) \cong g(\mathbf{x}^*) + \nabla_{\mathbf{x}} g(\mathbf{x}^*)^T (\mathbf{x} - \mathbf{x}^*) + \frac{1}{2} (\mathbf{x} - \mathbf{x}^*)^T \nabla_{\mathbf{x}}^2 g(\mathbf{x}^*) (\mathbf{x} - \mathbf{x}^*) \quad (1.27)$$

4. same as AMV+/AMV²+, except that the expansions are performed in $\mathbf{u}$ -space. The first-order option has been termed the $\mathbf{u}$ -space AMV+ method.

$$G(\mathbf{u}) \cong G(\mathbf{u}^*) + \nabla_{\mathbf{u}} G(\mathbf{u}^*)^T (\mathbf{u} - \mathbf{u}^*) \quad (1.28)$$

and the second-order option has been named the $\mathbf{u}$ -space AMV²+ method:

$$G(\mathbf{u}) \cong G(\mathbf{u}^*) + \nabla_{\mathbf{u}} G(\mathbf{u}^*)^T (\mathbf{u} - \mathbf{u}^*) + \frac{1}{2} (\mathbf{u} - \mathbf{u}^*)^T \nabla_{\mathbf{u}}^2 G(\mathbf{u}^*) (\mathbf{u} - \mathbf{u}^*) \quad (1.29)$$

5. a multipoint approximation in $\mathbf{x}$ -space. This approach involves a Taylor series approximation in intermediate variables where the powers used for the intermediate variables are selected to match information at the

current and previous expansion points. Based on the two-point exponential approximation concept (TPEA, [41]), the two-point adaptive nonlinearity approximation (TANA-3, [106]) approximates the limit state as:

$$g(\mathbf{x}) \cong g(\mathbf{x}_2) + \sum_{i=1}^n \frac{\partial g}{\partial x_i}(\mathbf{x}_2) \frac{x_{i,2}^{1-p_i}}{p_i} (x_i^{p_i} - x_{i,2}^{p_i}) + \frac{1}{2} \epsilon(\mathbf{x}) \sum_{i=1}^n (x_i^{p_i} - x_{i,2}^{p_i})^2 \quad (1.30)$$

where n is the number of uncertain variables and:

$$p_i = 1 + \ln \left[\frac{\frac{\partial g}{\partial x_i}(\mathbf{x}_1)}{\frac{\partial g}{\partial x_i}(\mathbf{x}_2)} \right] \bigg/ \ln \left[\frac{x_{i,1}}{x_{i,2}} \right] \quad (1.31)$$

$$\epsilon(\mathbf{x}) = \frac{H}{\sum_{i=1}^n (x_i^{p_i} - x_{i,1}^{p_i})^2 + \sum_{i=1}^n (x_i^{p_i} - x_{i,2}^{p_i})^2} \quad (1.32)$$

$$H = 2 \left[g(\mathbf{x}_1) - g(\mathbf{x}_2) - \sum_{i=1}^n \frac{\partial g}{\partial x_i}(\mathbf{x}_2) \frac{x_{i,2}^{1-p_i}}{p_i} (x_{i,1}^{p_i} - x_{i,2}^{p_i}) \right] \quad (1.33)$$

and $\mathbf{x}_2$ and $\mathbf{x}_1$ are the current and previous MPP estimates in x-space, respectively. Prior to the availability of two MPP estimates, x-space AMV+ is used.

6. a multipoint approximation in u-space. The u-space TANA-3 approximates the limit state as:

$$G(\mathbf{u}) \cong G(\mathbf{u}_2) + \sum_{i=1}^n \frac{\partial G}{\partial u_i}(\mathbf{u}_2) \frac{u_{i,2}^{1-p_i}}{p_i} (u_i^{p_i} - u_{i,2}^{p_i}) + \frac{1}{2} \epsilon(\mathbf{u}) \sum_{i=1}^n (u_i^{p_i} - u_{i,2}^{p_i})^2 \quad (1.34)$$

where:

$$p_i = 1 + \ln \left[\frac{\frac{\partial G}{\partial u_i}(\mathbf{u}_1)}{\frac{\partial G}{\partial u_i}(\mathbf{u}_2)} \right] \bigg/ \ln \left[\frac{u_{i,1}}{u_{i,2}} \right] \quad (1.35)$$

$$\epsilon(\mathbf{u}) = \frac{H}{\sum_{i=1}^n (u_i^{p_i} - u_{i,1}^{p_i})^2 + \sum_{i=1}^n (u_i^{p_i} - u_{i,2}^{p_i})^2} \quad (1.36)$$

$$H = 2 \left[G(\mathbf{u}_1) - G(\mathbf{u}_2) - \sum_{i=1}^n \frac{\partial G}{\partial u_i}(\mathbf{u}_2) \frac{u_{i,2}^{1-p_i}}{p_i} (u_{i,1}^{p_i} - u_{i,2}^{p_i}) \right] \quad (1.37)$$

and $\mathbf{u}_2$ and $\mathbf{u}_1$ are the current and previous MPP estimates in u-space, respectively. Prior to the availability of two MPP estimates, u-space AMV+ is used.

7. the MPP search on the original response functions without the use of any approximations. Combining this option with first-order and second-order integration approaches (see next section) results in the traditional first-order and second-order reliability methods (FORM and SORM).

The Hessian matrices in AMV² and AMV²+ may be available analytically, estimated numerically, or approximated through quasi-Newton updates. The selection between x-space or u-space for performing approximations depends on where the approximation will be more accurate, since this will result in more accurate MPP estimates (AMV, AMV²) or faster convergence (AMV+, AMV²+, TANA). Since this relative accuracy depends on the forms of the limit state $g(x)$ and the transformation $T(x)$ and is therefore application dependent in general, Dakota supports both options. A concern with approximation-based iterative search methods (i.e., AMV+, AMV²+ and TANA) is the robustness of their convergence to the MPP. It is possible for the MPP iterates to oscillate or even diverge. However, to date, this occurrence has been relatively rare, and Dakota contains checks

that monitor for this behavior. Another concern with TANA is numerical safeguarding (e.g., the possibility of raising negative x_i or u_i values to nonintegral p_i exponents in Equations 1.30, 1.32-1.34, and 1.36-1.37). Safe-guarding involves offsetting negative x_i or u_i and, for potential numerical difficulties with the logarithm ratios in Equations 1.31 and 1.35, reverting to either the linear ($p_i = 1$) or reciprocal ($p_i = -1$) approximation based on which approximation has lower error in $\frac{\partial g}{\partial x_i}(\mathbf{x}_1)$ or $\frac{\partial G}{\partial u_i}(\mathbf{u}_1)$.

1.1.2.2 Probability integrations

The second algorithmic variation involves the integration approach for computing probabilities at the MPP, which can be selected to be first-order (Equations 1.6-1.7) or second-order integration. Second-order integration involves applying a curvature correction [14, 57, 58]. Breitung applies a correction based on asymptotic analysis [14]:

$$p = \Phi(-\beta_p) \prod_{i=1}^{n-1} \frac{1}{\sqrt{1 + \beta_p \kappa_i}} \quad (1.38)$$

where κ_i are the principal curvatures of the limit state function (the eigenvalues of an orthonormal transformation of $\nabla_{\mathbf{u}}^2 G$, taken positive for a convex limit state) and $\beta_p \geq 0$ (a CDF or CCDF probability correction is selected to obtain the correct sign for β_p). An alternate correction in [57] is consistent in the asymptotic regime ($\beta_p \rightarrow \infty$) but does not collapse to first-order integration for $\beta_p = 0$:

$$p = \Phi(-\beta_p) \prod_{i=1}^{n-1} \frac{1}{\sqrt{1 + \psi(-\beta_p) \kappa_i}} \quad (1.39)$$

where $\psi() = \frac{\phi()}{\Phi()}$ and $\phi()$ is the standard normal density function. [58] applies further corrections to Equation 1.39 based on point concentration methods. At this time, all three approaches are available within the code, but the Hohenbichler-Rackwitz correction is used by default (switching the correction is a compile-time option in the source code and has not been exposed in the input specification).

1.1.2.3 Hessian approximations

To use a second-order Taylor series or a second-order integration when second-order information ($\nabla_{\mathbf{x}}^2 g$, $\nabla_{\mathbf{u}}^2 G$, and/or κ) is not directly available, one can estimate the missing information using finite differences or approximate it through use of quasi-Newton approximations. These procedures will often be needed to make second-order approaches practical for engineering applications.

In the finite difference case, numerical Hessians are commonly computed using either first-order forward differences of gradients using

$$\nabla^2 g(\mathbf{x}) \cong \frac{\nabla g(\mathbf{x} + h\mathbf{e}_i) - \nabla g(\mathbf{x})}{h} \quad (1.40)$$

to estimate the i^{th} Hessian column when gradients are analytically available, or second-order differences of function values using

$$\nabla^2 g(\mathbf{x}) \cong \frac{g(\mathbf{x} + h\mathbf{e}_i + h\mathbf{e}_j) - g(\mathbf{x} + h\mathbf{e}_i - h\mathbf{e}_j) - g(\mathbf{x} - h\mathbf{e}_i + h\mathbf{e}_j) + g(\mathbf{x} - h\mathbf{e}_i - h\mathbf{e}_j)}{4h^2} \quad (1.41)$$

to estimate the ij^{th} Hessian term when gradients are not directly available. This approach has the advantage of locally-accurate Hessians for each point of interest (which can lead to quadratic convergence rates in discrete Newton methods), but has the disadvantage that numerically estimating each of the matrix terms can be expensive.

Quasi-Newton approximations, on the other hand, do not reevaluate all of the second-order information for every point of interest. Rather, they accumulate approximate curvature information over time using secant updates. Since they utilize the existing gradient evaluations, they do not require any additional function evaluations for evaluating the Hessian terms. The quasi-Newton approximations of interest include the Broyden-Fletcher-Goldfarb-Shanno (BFGS) update

$$\mathbf{B}_{k+1} = \mathbf{B}_k - \frac{\mathbf{B}_k \mathbf{s}_k \mathbf{s}_k^T \mathbf{B}_k}{\mathbf{s}_k^T \mathbf{B}_k \mathbf{s}_k} + \frac{\mathbf{y}_k \mathbf{y}_k^T}{\mathbf{y}_k^T \mathbf{s}_k} \quad (1.42)$$

which yields a sequence of symmetric positive definite Hessian approximations, and the Symmetric Rank 1 (SR1) update

$$\mathbf{B}_{k+1} = \mathbf{B}_k + \frac{(\mathbf{y}_k - \mathbf{B}_k \mathbf{s}_k)(\mathbf{y}_k - \mathbf{B}_k \mathbf{s}_k)^T}{(\mathbf{y}_k - \mathbf{B}_k \mathbf{s}_k)^T \mathbf{s}_k} \quad (1.43)$$

which yields a sequence of symmetric, potentially indefinite, Hessian approximations. $\mathbf{B}_k$ is the k^{th} approximation to the Hessian $\nabla^2 g$, $\mathbf{s}_k = \mathbf{x}_{k+1} - \mathbf{x}_k$ is the step and $\mathbf{y}_k = \nabla g_{k+1} - \nabla g_k$ is the corresponding yield in the gradients. The selection of BFGS versus SR1 involves the importance of retaining positive definiteness in the Hessian approximations; if the procedure does not require it, then the SR1 update can be more accurate if the true Hessian is not positive definite. Initial scalings for $\mathbf{B}_0$ and numerical safeguarding techniques (damped BFGS, update skipping) are described in [34].

1.1.2.4 Optimization algorithms

The next algorithmic variation involves the optimization algorithm selection for solving Eqs. 1.15 and 1.16. The Hasofer-Lind Rackwitz-Fissler (HL-RF) algorithm [53] is a classical approach that has been broadly applied. It is a Newton-based approach lacking line search/trust region globalization, and is generally regarded as computationally efficient but occasionally unreliable. Dakota takes the approach of employing robust, general-purpose optimization algorithms with provable convergence properties. In particular, we employ the sequential quadratic programming (SQP) and nonlinear interior-point (NIP) optimization algorithms from the NPSOL [48] and OPT++ [69] libraries, respectively.

1.1.2.5 Warm Starting of MPP Searches

The final algorithmic variation for local reliability methods involves the use of warm starting approaches for improving computational efficiency. [33] describes the acceleration of MPP searches through warm starting with approximate iteration increment, with $z/p/\beta$ level increment, and with design variable increment. Warm started data includes the expansion point and associated response values and the MPP optimizer initial guess. Projections are used when an increment in $z/p/\beta$ level or design variables occurs. Warm starts were consistently effective in [33], with greater effectiveness for smaller parameter changes, and are used by default in Dakota.

1.2 Global Reliability Methods

Local reliability methods, while computationally efficient, have well-known failure mechanisms. When confronted with a limit state function that is nonsmooth, local gradient-based optimizers may stall due to gradient inaccuracy and fail to converge to an MPP. Moreover, if the limit state is multimodal (multiple MPPs), then a gradient-based local method can, at best, locate only one local MPP solution. Finally, a linear (Eqs. 1.6–1.7) or parabolic (Eqs. 1.38–1.39) approximation to the limit state at this MPP may fail to adequately capture the contour of a highly nonlinear limit state.

A reliability analysis method that is both efficient when applied to expensive response functions and accurate for a response function of any arbitrary shape is needed. This section develops such a method based on efficient global optimization [61] (EGO) to the search for multiple points on or near the limit state throughout the random variable space. By locating multiple points on the limit state, more complex limit states can be accurately modeled, resulting in a more accurate assessment of the reliability. It should be emphasized here that these multiple points exist on a single limit state. Because of its roots in efficient global optimization, this method of reliability analysis is called efficient global reliability analysis (EGRA) [?]. The following two subsections describe two capabilities that are incorporated into the EGRA algorithm: importance sampling and EGO.

1.2.1 Importance Sampling

An alternative to MPP search methods is to directly perform the probability integration numerically by sampling the response function. Sampling methods do not rely on a simplifying approximation to the shape of the limit state, so they can be more accurate than FORM and SORM, but they can also be prohibitively expensive because they generally require a large number of response function evaluations. Importance sampling methods reduce this expense by focusing the samples in the important regions of the uncertain space. They do this by centering the sampling density function at the MPP rather than at the mean. This ensures the samples will lie the region of interest, thus increasing the efficiency of the sampling method. Adaptive importance sampling (AIS) further improves the efficiency by adaptively updating the sampling density function. Multimodal adaptive importance sampling [27, 109] is a variation of AIS that allows for the use of multiple sampling densities making it better suited for cases where multiple sections of the limit state are highly probable.

Note that importance sampling methods require that the location of at least one MPP be known because it is used to center the initial sampling density. However, current gradient-based, local search methods used in MPP search may fail to converge or may converge to poor solutions for highly nonlinear problems, possibly making these methods inapplicable. As the next section describes, EGO is a global optimization method that does not depend on the availability of accurate gradient information, making convergence more reliable for nonsmooth response functions. Moreover, EGO has the ability to locate multiple failure points, which would provide multiple starting points and thus a good multimodal sampling density for the initial steps of multimodal AIS. The resulting Gaussian process model is accurate in the vicinity of the limit state, thereby providing an inexpensive surrogate that can be used to provide response function samples. As will be seen, using EGO to locate multiple points along the limit state, and then using the resulting Gaussian process model to provide function evaluations in multimodal AIS for the probability integration, results in an accurate and efficient reliability analysis tool.

1.2.2 Efficient Global Optimization

Efficient Global Optimization (EGO) was developed to facilitate the unconstrained minimization of expensive implicit response functions. The method builds an initial Gaussian process model as a global surrogate for the response function, then intelligently selects additional samples to be added for inclusion in a new Gaussian process model in subsequent iterations. The new samples are selected based on how much they are expected to improve the current best solution to the optimization problem. When this expected improvement is acceptably small, the globally optimal solution has been found. The application of this methodology to equality-constrained reliability analysis is the primary contribution of EGRA.

Efficient global optimization was originally proposed by Jones et al. [61] and has been adapted into similar methods such as sequential kriging optimization (SKO) [60]. The main difference between SKO and EGO lies within the specific formulation of what is known as the expected improvement function (EIF), which is the feature that sets all EGO/SKO-type methods apart from other global optimization methods. The EIF is used to select the location at which a new training point should be added to the Gaussian process model by maximizing the amount

of improvement in the objective function that can be expected by adding that point. A point could be expected to produce an improvement in the objective function if its predicted value is better than the current best solution, or if the uncertainty in its prediction is such that the probability of it producing a better solution is high. Because the uncertainty is higher in regions of the design space with fewer observations, this provides a balance between exploiting areas of the design space that predict good solutions, and exploring areas where more information is needed.

The general procedure of these EGO-type methods is:

1. Build an initial Gaussian process model of the objective function.
2. Find the point that maximizes the EIF. If the EIF value at this point is sufficiently small, stop.
3. Evaluate the objective function at the point where the EIF is maximized. Update the Gaussian process model using this new point. Go to Step 2.

The following sections discuss the construction of the Gaussian process model used, the form of the EIF, and then a description of how that EIF is modified for application to reliability analysis.

1.2.2.1 Gaussian Process Model

Gaussian process (GP) models are set apart from other surrogate models because they provide not just a predicted value at an unsampled point, but also an estimate of the prediction variance. This variance gives an indication of the uncertainty in the GP model, which results from the construction of the covariance function. This function is based on the idea that when input points are near one another, the correlation between their corresponding outputs will be high. As a result, the uncertainty associated with the model's predictions will be small for input points which are near the points used to train the model, and will increase as one moves further from the training points.

It is assumed that the true response function being modeled $G(\mathbf{u})$ can be described by: [23]

$$G(\mathbf{u}) = \mathbf{h}(\mathbf{u})^T \boldsymbol{\beta} + Z(\mathbf{u}) \quad (1.44)$$

where $\mathbf{h}(\cdot)$ is the trend of the model, $\boldsymbol{\beta}$ is the vector of trend coefficients, and $Z(\cdot)$ is a stationary Gaussian process with zero mean (and covariance defined below) that describes the departure of the model from its underlying trend. The trend of the model can be assumed to be any function, but taking it to be a constant value has been reported to be generally sufficient. [86] For the work presented here, the trend is assumed constant and $\boldsymbol{\beta}$ is taken as simply the mean of the responses at the training points. The covariance between outputs of the Gaussian process $Z(\cdot)$ at points $\mathbf{a}$ and $\mathbf{b}$ is defined as:

$$\text{Cov}[Z(\mathbf{a}), Z(\mathbf{b})] = \sigma_Z^2 R(\mathbf{a}, \mathbf{b}) \quad (1.45)$$

where σ_Z^2 is the process variance and $R(\cdot)$ is the correlation function. There are several options for the correlation function, but the squared-exponential function is common [86], and is used here for $R(\cdot)$:

$$R(\mathbf{a}, \mathbf{b}) = \exp \left[- \sum_{i=1}^d \theta_i (a_i - b_i)^2 \right] \quad (1.46)$$

where d represents the dimensionality of the problem (the number of random variables), and θ_i is a scale parameter that indicates the correlation between the points within dimension i . A large θ_i is representative of a short correlation length.

The expected value $\mu_G()$ and variance $\sigma_G^2()$ of the GP model prediction at point $\mathbf{u}$ are:

$$\mu_G(\mathbf{u}) = \mathbf{h}(\mathbf{u})^T \boldsymbol{\beta} + \mathbf{r}(\mathbf{u})^T \mathbf{R}^{-1}(\mathbf{g} - \mathbf{F}\boldsymbol{\beta}) \quad (1.47)$$

$$\sigma_G^2(\mathbf{u}) = \sigma_Z^2 - \begin{bmatrix} \mathbf{h}(\mathbf{u})^T & \mathbf{r}(\mathbf{u})^T \end{bmatrix} \begin{bmatrix} \mathbf{0} & \mathbf{F}^T \\ \mathbf{F} & \mathbf{R} \end{bmatrix}^{-1} \begin{bmatrix} \mathbf{h}(\mathbf{u}) \\ \mathbf{r}(\mathbf{u}) \end{bmatrix} \quad (1.48)$$

where $\mathbf{r}(\mathbf{u})$ is a vector containing the covariance between $\mathbf{u}$ and each of the n training points (defined by Eq. 1.45), $\mathbf{R}$ is an $n \times n$ matrix containing the correlation between each pair of training points, $\mathbf{g}$ is the vector of response outputs at each of the training points, and $\mathbf{F}$ is an $n \times q$ matrix with rows $\mathbf{h}(\mathbf{u}_i)^T$ (the trend function for training point i containing q terms; for a constant trend $q=1$). This form of the variance accounts for the uncertainty in the trend coefficients $\boldsymbol{\beta}$, but assumes that the parameters governing the covariance function (σ_Z^2 and $\boldsymbol{\theta}$) have known values.

The parameters σ_Z^2 and $\boldsymbol{\theta}$ are determined through maximum likelihood estimation. This involves taking the log of the probability of observing the response values $\mathbf{g}$ given the covariance matrix $\mathbf{R}$, which can be written as: [86]

$$\log [p(\mathbf{g}|\mathbf{R})] = -\frac{1}{n} \log |\mathbf{R}| - \log(\hat{\sigma}_Z^2) \quad (1.49)$$

where $|\mathbf{R}|$ indicates the determinant of $\mathbf{R}$, and $\hat{\sigma}_Z^2$ is the optimal value of the variance given an estimate of $\boldsymbol{\theta}$ and is defined by:

$$\hat{\sigma}_Z^2 = \frac{1}{n} (\mathbf{g} - \mathbf{F}\boldsymbol{\beta})^T \mathbf{R}^{-1} (\mathbf{g} - \mathbf{F}\boldsymbol{\beta}) \quad (1.50)$$

Maximizing Eq. 1.49 gives the maximum likelihood estimate of $\boldsymbol{\theta}$, which in turn defines σ_Z^2 .

1.2.2.2 Expected Improvement Function

The expected improvement function is used to select the location at which a new training point should be added. The EIF is defined as the expectation that any point in the search space will provide a better solution than the current best solution based on the expected values and variances predicted by the GP model. An important feature of the EIF is that it provides a balance between exploiting areas of the design space where good solutions have been found, and exploring areas of the design space where the uncertainty is high. First, recognize that at any point in the design space, the GP prediction $\hat{G}()$ is a Gaussian distribution:

$$\hat{G}(\mathbf{u}) \sim N[\mu_G(\mathbf{u}), \sigma_G(\mathbf{u})] \quad (1.51)$$

where the mean $\mu_G()$ and the variance $\sigma_G^2()$ were defined in Eqs. 1.47 and 1.48, respectively. The EIF is defined as: [61]

$$EI(\hat{G}(\mathbf{u})) \equiv E \left[\max \left(G(\mathbf{u}^*) - \hat{G}(\mathbf{u}), 0 \right) \right] \quad (1.52)$$

where $G(\mathbf{u}^*)$ is the current best solution chosen from among the true function values at the training points (henceforth referred to as simply G^*). This expectation can then be computed by integrating over the distribution $\hat{G}(\mathbf{u})$ with G^* held constant:

$$EI(\hat{G}(\mathbf{u})) = \int_{-\infty}^{G^*} (G^* - G) \hat{G}(\mathbf{u}) dG \quad (1.53)$$

where G is a realization of $\hat{G}$. This integral can be expressed analytically as: [61]

$$EI(\hat{G}(\mathbf{u})) = (G^* - \mu_G) \Phi \left(\frac{G^* - \mu_G}{\sigma_G} \right) + \sigma_G \phi \left(\frac{G^* - \mu_G}{\sigma_G} \right) \quad (1.54)$$

where it is understood that μ_G and σ_G are functions of $\mathbf{u}$.

The point at which the EIF is maximized is selected as an additional training point. With the new training point added, a new GP model is built and then used to construct another EIF, which is then used to choose another new training point, and so on, until the value of the EIF at its maximized point is below some specified tolerance. In Ref. [60] this maximization is performed using a Nelder-Mead simplex approach, which is a local optimization method. Because the EIF is often highly multimodal [61] it is expected that Nelder-Mead may fail to converge to the true global optimum. In Ref. [61], a branch-and-bound technique for maximizing the EIF is used, but was found to often be too expensive to run to convergence. In Dakota, an implementation of the DIRECT global optimization algorithm is used [44].

It is important to understand how the use of this EIF leads to optimal solutions. Eq. 1.54 indicates how much the objective function value at $\mathbf{x}$ is expected to be less than the predicted value at the current best solution. Because the GP model provides a Gaussian distribution at each predicted point, expectations can be calculated. Points with good expected values and even a small variance will have a significant expectation of producing a better solution (exploitation), but so will points that have relatively poor expected values and greater variance (exploration).

The application of EGO to reliability analysis, however, is made more complicated due to the inclusion of equality constraints (see Eqs. 1.15-1.16). For inverse reliability analysis, this extra complication is small. The response being modeled by the GP is the objective function of the optimization problem (see Eq. 1.16) and the deterministic constraint might be handled through the use of a merit function, thereby allowing EGO to solve this equality-constrained optimization problem. Here the problem lies in the interpretation of the constraint for multimodal problems as mentioned previously. In the forward reliability case, the response function appears in the constraint rather than the objective. Here, the maximization of the EIF is inappropriate because feasibility is the main concern. This application is therefore a significant departure from the original objective of EGO and requires a new formulation. For this problem, the expected feasibility function is introduced.

1.2.2.3 Expected Feasibility Function

The expected improvement function provides an indication of how much the true value of the response at a point can be expected to be less than the current best solution. It therefore makes little sense to apply this to the forward reliability problem where the goal is not to minimize the response, but rather to find where it is equal to a specified threshold value. The expected feasibility function (EFF) is introduced here to provide an indication of how well the true value of the response is expected to satisfy the equality constraint $G(\mathbf{u}) = \bar{z}$. Inspired by the contour estimation work in [81], this expectation can be calculated in a similar fashion as Eq. 1.53 by integrating over a region in the immediate vicinity of the threshold value $\bar{z} \pm \epsilon$:

$$EF(\hat{G}(\mathbf{u})) = \int_{z-\epsilon}^{z+\epsilon} [\epsilon - |\bar{z} - G|] \hat{G}(\mathbf{u}) dG \quad (1.55)$$

where G denotes a realization of the distribution $\hat{G}$, as before. Allowing z^+ and z^- to denote $\bar{z} \pm \epsilon$, respectively, this integral can be expressed analytically as:

$$\begin{aligned} EF(\hat{G}(\mathbf{u})) = & (\mu_G - \bar{z}) \left[2 \Phi \left(\frac{\bar{z} - \mu_G}{\sigma_G} \right) - \Phi \left(\frac{z^- - \mu_G}{\sigma_G} \right) - \Phi \left(\frac{z^+ - \mu_G}{\sigma_G} \right) \right] \\ & - \sigma_G \left[2 \phi \left(\frac{\bar{z} - \mu_G}{\sigma_G} \right) - \phi \left(\frac{z^- - \mu_G}{\sigma_G} \right) - \phi \left(\frac{z^+ - \mu_G}{\sigma_G} \right) \right] \\ & + \epsilon \left[\Phi \left(\frac{z^+ - \mu_G}{\sigma_G} \right) - \Phi \left(\frac{z^- - \mu_G}{\sigma_G} \right) \right] \end{aligned} \quad (1.56)$$

where ϵ is proportional to the standard deviation of the GP predictor ($\epsilon \propto \sigma_G$). In this case, z^- , z^+ , μ_G , σ_G , and ϵ are all functions of the location $\mathbf{u}$, while $\bar{z}$ is a constant. Note that the EFF provides the same balance between

exploration and exploitation as is captured in the EIF. Points where the expected value is close to the threshold ($\mu_G \approx \bar{z}$) and points with a large uncertainty in the prediction will have large expected feasibility values.

Chapter 2

Stochastic Expansion Methods

This chapter explores two approaches to forming stochastic expansions, the polynomial chaos expansion (PCE), which employs bases of multivariate orthogonal polynomials, and stochastic collocation (SC), which employs bases of multivariate interpolation polynomials. Both approaches capture the functional relationship between a set of output response metrics and a set of input random variables.

2.1 Orthogonal polynomials

2.1.1 Askey scheme

Table 2.1 shows the set of classical orthogonal polynomials which provide an optimal basis for different continuous probability distribution types. It is derived from the family of hypergeometric orthogonal polynomials known as the Askey scheme [7], for which the Hermite polynomials originally employed by Wiener [98] are a subset. The optimality of these basis selections derives from their orthogonality with respect to weighting functions that correspond to the probability density functions (PDFs) of the continuous distributions when placed in a standard form. The density and weighting functions differ by a constant factor due to the requirement that the integral of the PDF over the support range is one.

Table 2.1: Linkage between standard forms of continuous probability distributions and Askey scheme of continuous hyper-geometric polynomials.

Distribution	Density function	Polynomial	Weight function	Support range
Normal	$\frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}}$	Hermite $He_n(x)$	$e^{-\frac{x^2}{2}}$	$[-\infty, \infty]$
Uniform	$\frac{1}{2}$	Legendre $P_n(x)$	1	$[-1, 1]$
Beta	$\frac{(1-x)^\alpha (1+x)^\beta}{2^{\alpha+\beta+1} B(\alpha+1, \beta+1)}$	Jacobi $P_n^{(\alpha, \beta)}(x)$	$(1-x)^\alpha (1+x)^\beta$	$[-1, 1]$
Exponential	e^{-x}	Laguerre $L_n(x)$	e^{-x}	$[0, \infty]$
Gamma	$\frac{x^\alpha e^{-x}}{\Gamma(\alpha+1)}$	Generalized Laguerre $L_n^{(\alpha)}(x)$	$x^\alpha e^{-x}$	$[0, \infty]$

Note that Legendre is a special case of Jacobi for $\alpha = \beta = 0$, Laguerre is a special case of generalized Laguerre for $\alpha = 0$, $\Gamma(a)$ is the Gamma function which extends the factorial function to continuous values, and $B(a, b)$ is the Beta function defined as $B(a, b) = \frac{\Gamma(a)\Gamma(b)}{\Gamma(a+b)}$. Some care is necessary when specifying the α and β parameters

for the Jacobi and generalized Laguerre polynomials since the orthogonal polynomial conventions [1] differ from the common statistical PDF conventions. The former conventions are used in Table 2.1.

2.1.2 Numerically generated orthogonal polynomials

If all random inputs can be described using independent normal, uniform, exponential, beta, and gamma distributions, then Askey polynomials can be directly applied. If correlation or other distribution types are present, then additional techniques are required. One solution is to employ nonlinear variable transformations as described in Section 2.5 such that an Askey basis can be applied in the transformed space. This can be effective as shown in [39], but convergence rates are typically degraded. In addition, correlation coefficients are warped by the nonlinear transformation [26], and simple expressions for these transformed correlation values are not always readily available. An alternative is to numerically generate the orthogonal polynomials (using Gauss-Wigert [87], discretized Stieltjes [45], Chebyshev [45], or Gramm-Schmidt [99] approaches) and then compute their Gauss points and weights (using the Golub-Welsch [52] tridiagonal eigensolution). These solutions are optimal for given random variable sets having arbitrary probability density functions and eliminate the need to induce additional nonlinearity through variable transformations, but performing this process for general joint density functions with correlation is a topic of ongoing research (refer to Section 2.5 for additional details).

2.2 Interpolation polynomials

Interpolation polynomials may be either local or global and either value-based or gradient-enhanced: Lagrange (global value-based), Hermite (global gradient-enhanced), piecewise linear spline (local value-based), and piecewise cubic spline (local gradient-enhanced). Each of these combinations can be used within nodal or hierarchical interpolation formulations. The subsections that follow describe the one-dimensional interpolation polynomials for these cases and Section 2.4 describes their use for multivariate interpolation within the stochastic collocation algorithm.

2.2.1 Nodal interpolation

For value-based interpolation of a response function R in one dimension at an interpolation level l containing m^l points, the expression

$$R(\xi) \cong I^l(R) = \sum_{j=1}^{m_l} r(\xi_j) L_j(\xi) \quad (2.1)$$

reproduces the response values $r(\xi_j)$ at the interpolation points and smoothly interpolates between these values at other points. As we refine the interpolation level, we increase the number of collocation points in the rule and the number of interpolated response values.

For the case of gradient-enhancement, interpolation of a one-dimensional function involves both type 1 and type 2 interpolation polynomials,

$$R(\xi) \cong I^l(R) = \sum_{j=1}^{m_l} \left[r(\xi_j) H_j^{(1)}(\xi) + \frac{dr}{d\xi}(\xi_j) H_j^{(2)}(\xi) \right] \quad (2.2)$$

where the former interpolate a particular value while producing a zero gradient (i^{th} type 1 interpolant produces a value of 1 for the i^{th} collocation point, zero values for all other points, and zero gradients for all points) and the

latter interpolate a particular gradient while producing a zero value (i^{th} type 2 interpolant produces a gradient of 1 for the i^{th} collocation point, zero gradients for all other points, and zero values for all points).

2.2.1.1 Global value-based

Lagrange polynomials interpolate a set of points in a single dimension using the functional form

$$L_j = \prod_{\substack{k=1 \\ k \neq j}}^m \frac{\xi - \xi_k}{\xi_j - \xi_k} \quad (2.3)$$

where it is evident that L_j is 1 at $\xi = \xi_j$, is 0 for each of the points $\xi = \xi_k$, and has order $m - 1$.

To improve numerical efficiency and stability, a barycentric Lagrange formulation [10, 55] is used. We define the barycentric weights w_j as

$$w_j = \prod_{\substack{k=1 \\ k \neq j}}^m \frac{1}{\xi_j - \xi_k} \quad (2.4)$$

and we precompute them for a given interpolation point set $\xi_j, j \in 1, \dots, m$. Then, defining the quantity $l(\xi)$ as

$$l(\xi) = \prod_{k=1}^m (\xi - \xi_k) \quad (2.5)$$

which will be computed for each new interpolated point ξ , we can rewrite Eq. 2.1 as

$$R(\xi) = l(\xi) \sum_{j=1}^m \frac{w_j}{x - x_j} r(\xi_j) \quad (2.6)$$

where much of the computational work has been moved outside the summation. Eq. 2.6 is the first form of barycentric interpolation. Using an identity from the interpolation of unity ($R(\xi) = 1$ and each $r(\xi_j) = 1$ in Eq. 2.6) to eliminate $l(x)$, we arrive at the second form of the barycentric interpolation formula:

$$R(\xi) = \frac{\sum_{j=1}^m \frac{w_j}{x - x_j} r(\xi_j)}{\sum_{j=1}^m \frac{w_j}{x - x_j}} \quad (2.7)$$

For both formulations, we reduce the computational effort for evaluating the interpolant from $O(m^2)$ to $O(m)$ operations per interpolated point, with the penalty of requiring additional care to avoid division by zero when ξ matches one of the ξ_j . Relative to the first form, the second form has the additional advantage that common factors within the w_j can be canceled (possible for Clenshaw-Curtis and Newton-Cotes point sets, but not for general Gauss points), further reducing the computational requirements. Barycentric formulations can also be used for hierarchical interpolation (Section 2.2.2) with Lagrange interpolation polynomials, but they are not applicable to local spline or gradient-enhanced Hermite interpolants.

2.2.1.2 Global gradient-enhanced

Hermite interpolation polynomials (not to be confused with Hermite orthogonal polynomials shown in Table 2.1) interpolate both values and derivatives. In our case, we are interested in interpolating values and first derivatives, i.e., gradients. One-dimensional polynomials satisfying the interpolation constraints for general point sets are generated using divided differences as described in [16].

2.2.1.3 Local value-based

Linear spline basis polynomials define a “hat function,” which produces the value of one at its collocation point and decays linearly to zero at its nearest neighbors. In the case where its collocation point corresponds to a domain boundary, then the half interval that extends beyond the boundary is truncated.

For the case of non-equidistant closed points (e.g., Clenshaw-Curtis), the linear spline polynomials are defined as

$$L_j(\xi) = \begin{cases} 1 - \frac{\xi - \xi_j}{\xi_{j+1} - \xi_j} & \text{if } \xi_{j-1} \leq \xi \leq \xi_j \text{ (left half interval)} \\ 1 - \frac{\xi - \xi_j}{\xi_{j+1} - \xi_j} & \text{if } \xi_j < \xi \leq \xi_{j+1} \text{ (right half interval)} \\ 0 & \text{otherwise} \end{cases} \quad (2.8)$$

For the case of equidistant closed points (i.e., Newton-Cotes), this can be simplified to

$$L_j(\xi) = \begin{cases} 1 - \frac{|\xi - \xi_j|}{h} & \text{if } |\xi - \xi_j| \leq h \\ 0 & \text{otherwise} \end{cases} \quad (2.9)$$

for h defining the half-interval $\frac{b-a}{m-1}$ of the hat function L_j over the range $\xi \in [a, b]$. For the special case of $m = 1$ point, $L_1(\xi) = 1$ for $\xi_1 = \frac{b+a}{2}$ in both cases above.

2.2.1.4 Local gradient-enhanced

Type 1 cubic spline interpolants are formulated as follows:

$$H_j^{(1)}(\xi) = \begin{cases} t^2(3 - 2t) & \text{for } t = \frac{\xi - \xi_{j-1}}{\xi_j - \xi_{j-1}} \text{ if } \xi_{j-1} \leq \xi \leq \xi_j \text{ (left half interval)} \\ (t - 1)^2(1 + 2t) & \text{for } t = \frac{\xi - \xi_j}{\xi_{j+1} - \xi_j} \text{ if } \xi_j < \xi \leq \xi_{j+1} \text{ (right half interval)} \\ 0 & \text{otherwise} \end{cases} \quad (2.10)$$

which produce the desired zero-one-zero property for left-center-right values and zero-zero-zero property for left-center-right gradients. Type 2 cubic spline interpolants are formulated as follows:

$$H_j^{(2)}(\xi) = \begin{cases} ht^2(t - 1) & \text{for } h = \xi_j - \xi_{j-1}, \quad t = \frac{\xi - \xi_{j-1}}{h} \text{ if } \xi_{j-1} \leq \xi \leq \xi_j \text{ (left half interval)} \\ ht(t - 1)^2 & \text{for } h = \xi_{j+1} - \xi_j, \quad t = \frac{\xi - \xi_j}{h} \text{ if } \xi_j < \xi \leq \xi_{j+1} \text{ (right half interval)} \\ 0 & \text{otherwise} \end{cases} \quad (2.11)$$

which produce the desired zero-zero-zero property for left-center-right values and zero-one-zero property for left-center-right gradients. For the special case of $m = 1$ point over the range $\xi \in [a, b]$, $H_1^{(1)}(\xi) = 1$ and $H_1^{(2)}(\xi) = \xi$ for $\xi_1 = \frac{b+a}{2}$.

2.2.2 Hierarchical interpolation

In a hierarchical formulation, we reformulate the interpolation in terms of differences between interpolation levels:

$$\Delta^l(R) = I^l(R) - I^{l-1}(R), \quad l \geq 1 \quad (2.12)$$

where $I^l(R)$ is as defined in Eqs. 2.1–2.2 using the same local or global definitions for $L_j(\xi)$, $H_j^{(1)}(\xi)$, and $H_j^{(2)}(\xi)$, and $I^{l-1}(R)$ is evaluated as $I^l(I^{l-1}(R))$, indicating reinterpolation of the lower level interpolant across the higher level point set [63, 4].

Utilizing Eqs. 2.1–2.2, we can represent this difference interpolant as

$$\Delta^l(R) = \begin{cases} \sum_{j=1}^{m_l} [r(\xi_j) - I^{l-1}(R)(\xi_j)] L_j(\xi) & \text{value-based} \\ \sum_{j=1}^{m_l} \left([r(\xi_j) - I^{l-1}(R)(\xi_j)] H_j^{(1)}(\xi) + \left[\frac{dr}{d\xi}(\xi_j) - \frac{dI^{l-1}(R)}{d\xi}(\xi_j) \right] H_j^{(2)}(\xi) \right) & \text{gradient-enhanced} \end{cases} \quad (2.13)$$

where $I^{l-1}(R)(\xi_j)$ and $\frac{dI^{l-1}(R)}{d\xi}(\xi_j)$ are the value and gradient, respectively, of the lower level interpolant evaluated at the higher level points. We then define hierarchical surpluses $s, s^{(1)}, s^{(2)}$ at a point ξ_j as the bracketed terms in Eq 2.13. These surpluses can be interpreted as local interpolation error estimates since they capture the difference between the true values and the values predicted by the previous interpolant.

For the case where we use nested point sets among the interpolation levels, the interpolant differences for points contained in both sets are zero, allowing us to restrict the summations above to $\sum_{j=1}^{m_{\Delta_l}}$ where we define the set $\Xi_{\Delta_l} = \Xi_l \setminus \Xi_{l-1}$ that contains $m_{\Delta_l} = m_l - m_{l-1}$ points. $\Delta^l(R)$ then becomes

$$\Delta^l(R) = \begin{cases} \sum_{j=1}^{m_{\Delta_l}} s(\xi_j) L_j(\xi) & \text{value-based} \\ \sum_{j=1}^{m_{\Delta_l}} \left(s^{(1)}(\xi_j) H_j^{(1)}(\xi) + s^{(2)}(\xi_j) H_j^{(2)}(\xi) \right) & \text{gradient-enhanced} \end{cases} \quad (2.14)$$

The original interpolant $I^l(R)$ can be represented as a summation of these difference interpolants

$$I^l(R) = \Delta^l(R) + I^{l-1}(R) = \sum_{i=1}^l \Delta^i(R) \quad (2.15)$$

We will employ these hierarchical definitions within stochastic collocation on sparse grids in Section 2.4.3.

2.3 Generalized Polynomial Chaos

The set of polynomials from 2.1.1 and 2.1.2 are used as an orthogonal basis to approximate the functional form between the stochastic response output and each of its random inputs. The chaos expansion for a response R takes the form

$$R = a_0 B_0 + \sum_{i_1=1}^{\infty} a_{i_1} B_1(\xi_{i_1}) + \sum_{i_1=1}^{\infty} \sum_{i_2=1}^{i_1} a_{i_1 i_2} B_2(\xi_{i_1}, \xi_{i_2}) + \sum_{i_1=1}^{\infty} \sum_{i_2=1}^{i_1} \sum_{i_3=1}^{i_2} a_{i_1 i_2 i_3} B_3(\xi_{i_1}, \xi_{i_2}, \xi_{i_3}) + \dots \quad (2.16)$$

where the random vector dimension is unbounded and each additional set of nested summations indicates an additional order of polynomials in the expansion. This expression can be simplified by replacing the order-based indexing with a term-based indexing

$$R = \sum_{j=0}^{\infty} \alpha_j \Psi_j(\xi) \quad (2.17)$$

where there is a one-to-one correspondence between $a_{i_1 i_2 \dots i_n}$ and α_j and between $B_n(\xi_{i_1}, \xi_{i_2}, \dots, \xi_{i_n})$ and $\Psi_j(\xi)$. Each of the $\Psi_j(\xi)$ are multivariate polynomials which involve products of the one-dimensional polynomials. For example, a multivariate Hermite polynomial $B(\xi)$ of order n is defined from

$$B_n(\xi_{i_1}, \dots, \xi_{i_n}) = e^{\frac{1}{2}\xi^T \xi} (-1)^n \frac{\partial^n}{\partial \xi_{i_1} \dots \partial \xi_{i_n}} e^{-\frac{1}{2}\xi^T \xi} \quad (2.18)$$

which can be shown to be a product of one-dimensional Hermite polynomials involving an expansion term multi-index t_i^j :

$$B_n(\xi_{i_1}, \dots, \xi_{i_n}) = \Psi_j(\xi) = \prod_{i=1}^n \psi_{t_i^j}(\xi_i) \quad (2.19)$$

In the case of a mixed basis, the same multi-index definition is employed although the one-dimensional polynomials $\psi_{t_i^j}$ are heterogeneous in type.

2.3.1 Expansion truncation and tailoring

In practice, one truncates the infinite expansion at a finite number of random variables and a finite expansion order

$$R \cong \sum_{j=0}^P \alpha_j \Psi_j(\xi) \quad (2.20)$$

Traditionally, the polynomial chaos expansion includes a complete basis of polynomials up to a fixed total-order specification. That is, for an expansion of total order p involving n random variables, the expansion term multi-index defining the set of Ψ_j is constrained by

$$\sum_{i=1}^n t_i^j \leq p \quad (2.21)$$

For example, the multidimensional basis polynomials for a second-order expansion over two random dimensions are

$$\begin{aligned} \Psi_0(\xi) &= \psi_0(\xi_1) \psi_0(\xi_2) = 1 \\ \Psi_1(\xi) &= \psi_1(\xi_1) \psi_0(\xi_2) = \xi_1 \\ \Psi_2(\xi) &= \psi_0(\xi_1) \psi_1(\xi_2) = \xi_2 \\ \Psi_3(\xi) &= \psi_2(\xi_1) \psi_0(\xi_2) = \xi_1^2 - 1 \\ \Psi_4(\xi) &= \psi_1(\xi_1) \psi_1(\xi_2) = \xi_1 \xi_2 \\ \Psi_5(\xi) &= \psi_0(\xi_1) \psi_2(\xi_2) = \xi_2^2 - 1 \end{aligned}$$

The total number of terms N_t in an expansion of total order p involving n random variables is given by

$$N_t = 1 + P = 1 + \sum_{s=1}^p \frac{1}{s!} \prod_{r=0}^{s-1} (n+r) = \frac{(n+p)!}{n!p!} \quad (2.22)$$

This traditional approach will be referred to as a “total-order expansion.”

An important alternative approach is to employ a “tensor-product expansion,” in which polynomial order bounds are applied on a per-dimension basis (no total-order bound is enforced) and all combinations of the one-dimensional polynomials are included. That is, the expansion term multi-index defining the set of Ψ_j is constrained by

$$t_i^j \leq p_i \quad (2.23)$$

where p_i is the polynomial order bound for the i^{th} dimension. In this case, the example basis for $p = 2, n = 2$ is

$$\begin{aligned} \Psi_0(\xi) &= \psi_0(\xi_1) \psi_0(\xi_2) = 1 \\ \Psi_1(\xi) &= \psi_1(\xi_1) \psi_0(\xi_2) = \xi_1 \\ \Psi_2(\xi) &= \psi_2(\xi_1) \psi_0(\xi_2) = \xi_1^2 - 1 \\ \Psi_3(\xi) &= \psi_0(\xi_1) \psi_1(\xi_2) = \xi_2 \\ \Psi_4(\xi) &= \psi_1(\xi_1) \psi_1(\xi_2) = \xi_1 \xi_2 \\ \Psi_5(\xi) &= \psi_2(\xi_1) \psi_1(\xi_2) = (\xi_1^2 - 1) \xi_2 \\ \Psi_6(\xi) &= \psi_0(\xi_1) \psi_2(\xi_2) = \xi_2^2 - 1 \\ \Psi_7(\xi) &= \psi_1(\xi_1) \psi_2(\xi_2) = \xi_1 (\xi_2^2 - 1) \\ \Psi_8(\xi) &= \psi_2(\xi_1) \psi_2(\xi_2) = (\xi_1^2 - 1) (\xi_2^2 - 1) \end{aligned}$$

and the total number of terms N_t is

$$N_t = 1 + P = \prod_{i=1}^n (p_i + 1) \quad (2.24)$$

It is apparent from Eq. 2.24 that the tensor-product expansion readily supports anisotropy in polynomial order for each dimension, since the polynomial order bounds for each dimension can be specified independently. It is also feasible to support anisotropy with total-order expansions, through pruning polynomials that satisfy the total-order bound but violate individual per-dimension bounds (the number of these pruned polynomials would then be subtracted from Eq. 2.22). Finally, additional tailoring of the expansion form is used in the case of sparse grids (see Section 2.6.3) through the use of a summation of anisotropic tensor expansions. In all cases, the specifics of the expansion are codified in the term multi-index, and subsequent machinery for estimating response values and statistics from the expansion can be performed in a manner that is agnostic to the specific expansion form.

2.4 Stochastic Collocation

The SC expansion is formed as a sum of a set of multidimensional interpolation polynomials, one polynomial per interpolated response quantity (one response value and potentially multiple response gradient components) per unique collocation point.

2.4.1 Value-Based Nodal

For value-based interpolation in multiple dimensions, a tensor-product of the one-dimensional polynomials described in Section 2.2.1.1 or Section 2.2.1.3 is used:

$$R(\boldsymbol{\xi}) \cong \sum_{j_1=1}^{m_{i_1}} \cdots \sum_{j_n=1}^{m_{i_n}} r(\xi_{j_1}^{i_1}, \dots, \xi_{j_n}^{i_n}) (L_{j_1}^{i_1} \otimes \cdots \otimes L_{j_n}^{i_n}) \quad (2.25)$$

where $\mathbf{i} = (m_1, m_2, \dots, m_n)$ are the number of nodes used in the n -dimensional interpolation and $\xi_{j_k}^{i_k}$ indicates the j^{th} point out of i possible collocation points in the k^{th} dimension. This can be simplified to

$$R(\boldsymbol{\xi}) \cong \sum_{j=1}^{N_p} r_j L_j(\boldsymbol{\xi}) \quad (2.26)$$

where N_p is the number of unique collocation points in the multidimensional grid. The multidimensional interpolation polynomials are defined as

$$L_j(\boldsymbol{\xi}) = \prod_{k=1}^n L_{c_k^j}(\xi_k) \quad (2.27)$$

where c_k^j is a collocation multi-index (similar to the expansion term multi-index in Eq. 2.19) that maps from the j^{th} unique collocation point to the corresponding multidimensional indices within the tensor grid, and we have dropped the superscript notation indicating the number of nodes in each dimension for simplicity. The tensor-product structure preserves the desired interpolation properties where the j^{th} multivariate interpolation polynomial assumes the value of 1 at the j^{th} point and assumes the value of 0 at all other points, thereby reproducing the response values at each of the collocation points and smoothly interpolating between these values at other unsampled points. When the one-dimensional interpolation polynomials are defined using a barycentric formulation as described in Section 2.2.1.1 (i.e., Eq. 2.7), additional efficiency in evaluating a tensor interpolant

is achieved using the procedure in [64], which amounts to a multi-dimensional extension to Horner's rule for tensor-product polynomial evaluation.

Multivariate interpolation on Smolyak sparse grids involves a weighted sum of the tensor products in Eq. 2.25 with varying i levels. For sparse interpolants based on nested quadrature rules (e.g., Clenshaw-Curtis, Gauss-Patterson, Genz-Keister), the interpolation property is preserved, but sparse interpolants based on non-nested rules may exhibit some interpolation error at the collocation points.

2.4.2 Gradient-Enhanced Nodal

For gradient-enhanced interpolation in multiple dimensions, we extend the formulation in Eq. 2.26 to use a tensor-product of the one-dimensional type 1 and type 2 polynomials described in Section 2.2.1.2 or Section 2.2.1.4:

$$R(\boldsymbol{\xi}) \cong \sum_{j=1}^{N_p} \left[r_j \mathbf{H}_j^{(1)}(\boldsymbol{\xi}) + \sum_{k=1}^n \frac{dr_j}{d\xi_k} \mathbf{H}_{jk}^{(2)}(\boldsymbol{\xi}) \right] \quad (2.28)$$

The multidimensional type 1 basis polynomials are

$$\mathbf{H}_j^{(1)}(\boldsymbol{\xi}) = \prod_{k=1}^n H_{c_k^j}^{(1)}(\xi_k) \quad (2.29)$$

where c_k^j is the same collocation multi-index described for Eq. 2.27 and the superscript notation indicating the number of nodes in each dimension has again been omitted. The multidimensional type 2 basis polynomials for the k^{th} gradient component are the same as the type 1 polynomials for each dimension except k :

$$\mathbf{H}_{jk}^{(2)}(\boldsymbol{\xi}) = H_{c_k^j}^{(2)}(\xi_k) \prod_{\substack{l=1 \\ l \neq k}}^n H_{c_l^j}^{(1)}(\xi_l) \quad (2.30)$$

As for the value-based case, multivariate interpolation on Smolyak sparse grids involves a weighted sum of the tensor products in Eq. 2.28 with varying i levels.

2.4.3 Hierarchical

In the case of multivariate hierarchical interpolation on nested grids, we are interested in tensor products of the one-dimensional difference interpolants described in Section 2.2.2, with

$$\Delta^l(R) = \sum_{j_1=1}^{m_{\Delta_1}} \cdots \sum_{j_n=1}^{m_{\Delta_n}} s\left(\xi_{j_1}^{\Delta_1}, \dots, \xi_{j_n}^{\Delta_n}\right) \left(L_{j_1}^{\Delta_1} \otimes \cdots \otimes L_{j_n}^{\Delta_n}\right) \quad (2.31)$$

for value-based, and

$$\begin{aligned} \Delta^l(R) &= \sum_{j_1=1}^{m_{\Delta_1}} \cdots \sum_{j_n=1}^{m_{\Delta_n}} \left[s^{(1)}\left(\xi_{j_1}^{\Delta_1}, \dots, \xi_{j_n}^{\Delta_n}\right) \left(H_{j_1}^{(1)\Delta_1} \otimes \cdots \otimes H_{j_n}^{(1)\Delta_n}\right) + \sum_{k=1}^n s_k^{(2)}\left(\xi_{j_1}^{\Delta_1}, \dots, \xi_{j_n}^{\Delta_n}\right) \left(H_k^{(2)\Delta_1} \otimes \cdots \otimes H_k^{(2)\Delta_n}\right) \right] \end{aligned} \quad (2.32)$$

for gradient-enhanced, where k indicates the gradient component being interpolated.

These difference interpolants are particularly useful within sparse grid interpolation, for which the Δ^l can be employed directly within Eq. 2.40.

2.5 Transformations to uncorrelated standard variables

Polynomial chaos and stochastic collocation are expanded using polynomials that are functions of independent random variables ξ , which are often standardized forms of common distributions. Thus, a key component of stochastic expansion approaches is performing a transformation of variables from the original random variables $\mathbf{x}$ to independent (standard) random variables ξ and then applying the stochastic expansion in the transformed space. This notion of independent standard space is extended over the notion of “u-space” used in reliability methods (see Section 1.1.2) in that it extends the standardized set beyond standard normals. For distributions that are already independent, three different approaches are of interest:

1. *Extended basis:* For each Askey distribution type, employ the corresponding Askey basis (Table 2.1). For non-Askey types, numerically generate an optimal polynomial basis for each independent distribution as described in Section 2.1.2. These numerically-generated basis polynomials are not coerced into any standardized form, but rather employ the actual distribution parameters of the individual random variables. Thus, not even a linear variable transformation is employed for these variables. With usage of the optimal basis corresponding to each of the random variable types, we avoid inducing additional nonlinearity that can slow convergence.
2. *Askey basis:* For non-Askey types, perform a nonlinear variable transformation from a given input distribution to the most similar Askey basis. For example, lognormal distributions might employ a Hermite basis in a transformed standard normal space and loguniform, triangular, and histogram distributions might employ a Legendre basis in a transformed standard uniform space. All distributions then employ the Askey orthogonal polynomials and their associated Gauss points/weights.
3. *Wiener basis:* For non-normal distributions, employ a nonlinear variable transformation to standard normal distributions. All distributions then employ the Hermite orthogonal polynomials and their associated Gauss points/weights.

For dependent distributions, we must first perform a nonlinear variable transformation to uncorrelated standard normal distributions, due to the independence of decorrelated standard normals. This involves the Nataf transformation, described in the following paragraph. We then have the following choices:

1. *Single transformation:* Following the Nataf transformation to independent standard normal distributions, employ the Wiener basis in the transformed space.
2. *Double transformation:* From independent standard normal space, transform back to either the original marginal distributions or the desired Askey marginal distributions and employ an extended or Askey basis, respectively, in the transformed space. Independence is maintained, but the nonlinearity of the Nataf transformation is at least partially mitigated.

Dakota does not yet implement the double transformation concept, such that each correlated variable will employ a Wiener basis approach.

The transformation from correlated non-normal distributions to uncorrelated standard normal distributions is denoted as $\xi = T(\mathbf{x})$ with the reverse transformation denoted as $\mathbf{x} = T^{-1}(\xi)$. These transformations are nonlinear

in general, and possible approaches include the Rosenblatt [85], Nataf [26], and Box-Cox [12] transformations. Dakota employs the Nataf transformation, which is suitable for the common case when marginal distributions and a correlation matrix are provided, but full joint distributions are not known¹. The Nataf transformation occurs in the following two steps. To transform between the original correlated x-space variables and correlated standard normals (“z-space”), a CDF matching condition is applied for each of the marginal distributions:

$$\Phi(z_i) = F(x_i) \quad (2.33)$$

where $\Phi()$ is the standard normal cumulative distribution function and $F()$ is the cumulative distribution function of the original probability distribution. Then, to transform between correlated z-space variables and uncorrelated ξ -space variables, the Cholesky factor $\mathbf{L}$ of a modified correlation matrix is used:

$$\mathbf{z} = \mathbf{L}\boldsymbol{\xi} \quad (2.34)$$

where the original correlation matrix for non-normals in x-space has been modified to represent the corresponding “warped” correlation in z-space [26].

2.6 Spectral projection

The major practical difference between PCE and SC is that, in PCE, one must estimate the coefficients for known basis functions, whereas in SC, one must form the interpolants for known coefficients. PCE estimates its coefficients using either spectral projection or linear regression, where the former approach involves numerical integration based on random sampling, tensor-product quadrature, Smolyak sparse grids, or cubature methods. In SC, the multidimensional interpolants need to be formed over structured data sets, such as point sets from quadrature or sparse grids; approaches based on random sampling may not be used.

The spectral projection approach projects the response against each basis function using inner products and employs the polynomial orthogonality properties to extract each coefficient. Similar to a Galerkin projection, the residual error from the approximation is rendered orthogonal to the selected basis. From Eq. 2.20, taking the inner product of both sides with respect to Ψ_j and enforcing orthogonality yields:

$$\alpha_j = \frac{\langle R, \Psi_j \rangle}{\langle \Psi_j^2 \rangle} = \frac{1}{\langle \Psi_j^2 \rangle} \int_{\Omega} R \Psi_j \varrho(\boldsymbol{\xi}) d\boldsymbol{\xi}, \quad (2.35)$$

where each inner product involves a multidimensional integral over the support range of the weighting function. In particular, $\Omega = \Omega_1 \otimes \cdots \otimes \Omega_n$, with possibly unbounded intervals $\Omega_j \subset \mathbb{R}$ and the tensor product form $\varrho(\boldsymbol{\xi}) = \prod_{i=1}^n \varrho_i(\xi_i)$ of the joint probability density (weight) function. The denominator in Eq. 2.35 is the norm squared of the multivariate orthogonal polynomial, which can be computed analytically using the product of univariate norms squared

$$\langle \Psi_j^2 \rangle = \prod_{i=1}^n \langle \psi_{t_i^j}^2 \rangle \quad (2.36)$$

where the univariate inner products have simple closed form expressions for each polynomial in the Askey scheme [1] and are readily computed as part of the numerically-generated solution procedures described in Section 2.1.2. Thus, the primary computational effort resides in evaluating the numerator, which is evaluated numerically using sampling, quadrature, cubature, or sparse grid approaches (and this numerical approximation leads to use of the term “pseudo-spectral” by some investigators).

¹If joint distributions are known, then the Rosenblatt transformation is preferred.

2.6.1 Sampling

In the sampling approach, the integral evaluation is equivalent to computing the expectation (mean) of the response-basis function product (the numerator in Eq. 2.35) for each term in the expansion when sampling within the density of the weighting function. This approach is only valid for PCE and since sampling does not provide any particular monomial coverage guarantee, it is common to combine this coefficient estimation approach with a total-order chaos expansion.

In computational practice, coefficient estimations based on sampling benefit from first estimating the response mean (the first PCE coefficient) and then removing the mean from the expectation evaluations for all subsequent coefficients. While this has no effect for quadrature/sparse grid methods (see following two sections) and little effect for fully-resolved sampling, it does have a small but noticeable beneficial effect for under-resolved sampling.

2.6.2 Tensor product quadrature

In quadrature-based approaches, the simplest general technique for approximating multidimensional integrals, as in Eq. 2.35, is to employ a tensor product of one-dimensional quadrature rules. Since there is little benefit to the use of nested quadrature rules in the tensor-product case², we choose Gaussian abscissas, i.e. the zeros of polynomials that are orthogonal with respect to a density function weighting, e.g. Gauss-Hermite, Gauss-Legendre, Gauss-Laguerre, generalized Gauss-Laguerre, Gauss-Jacobi, or numerically-generated Gauss rules.

We first introduce an index $i \in \mathbb{N}_+$, $i \geq 1$. Then, for each value of i , let $\{\xi_1^i, \dots, \xi_{m_i}^i\} \subset \Omega_i$ be a sequence of abscissas for quadrature on Ω_i . For $f \in C^0(\Omega_i)$ and $n = 1$ we introduce a sequence of one-dimensional quadrature operators

$$\mathcal{W}^i(f)(\xi) = \sum_{j=1}^{m_i} f(\xi_j^i) w_j^i, \quad (2.37)$$

with $m_i \in \mathbb{N}$ given. When utilizing Gaussian quadrature, Eq. 2.37 integrates exactly all polynomials of degree less than $2m_i - 1$, for each $i = 1, \dots, n$. Given an expansion order p , the highest order coefficient evaluations (Eq. 2.35) can be assumed to involve integrands of at least polynomial order $2p$ (Ψ of order p and R modeled to order p) in each dimension such that a minimal Gaussian quadrature order of $p + 1$ will be required to obtain good accuracy in these coefficients.

Now, in the multivariate case $n > 1$, for each $f \in C^0(\Omega)$ and the multi-index $\mathbf{i} = (i_1, \dots, i_n) \in \mathbb{N}_+^n$ we define the full tensor product quadrature formulas

$$\mathcal{Q}_{\mathbf{i}}^n f(\xi) = (\mathcal{W}^{i_1} \otimes \dots \otimes \mathcal{W}^{i_n})(f)(\xi) = \sum_{j_1=1}^{m_{i_1}} \dots \sum_{j_n=1}^{m_{i_n}} f(\xi_{j_1}^{i_1}, \dots, \xi_{j_n}^{i_n}) (w_{j_1}^{i_1} \otimes \dots \otimes w_{j_n}^{i_n}). \quad (2.38)$$

Clearly, the above product needs $\prod_{j=1}^n m_{i_j}$ function evaluations. Therefore, when the number of input random variables is small, full tensor product quadrature is a very effective numerical tool. On the other hand, approximations based on tensor product grids suffer from the *curse of dimensionality* since the number of collocation points in a tensor grid grows exponentially fast in the number of input random variables. For example, if Eq. 2.38 employs the same order for all random dimensions, $m_{i_j} = m$, then Eq. 2.38 requires m^n function evaluations.

In [35], it is demonstrated that close synchronization of expansion form with the monomial resolution of a particular numerical integration technique can result in significant performance improvements. In particular, the traditional approach of employing a total-order PCE (Eqs. 2.21–2.22) neglects a significant portion of the monomial coverage for a tensor-product quadrature approach, and one should rather employ a tensor-product PCE

²Unless a refinement procedure is in use.

(Eqs. 2.23–2.24) to provide improved synchronization and more effective usage of the Gauss point evaluations. When the quadrature points are standard Gauss rules (i.e., no Clenshaw-Curtis, Gauss-Patterson, or Genz-Keister nested rules), it has been shown that tensor-product PCE and SC result in identical polynomial forms [22], completely eliminating a performance gap that exists between total-order PCE and SC [35].

2.6.3 Smolyak sparse grids

If the number of random variables is moderately large, one should rather consider sparse tensor product spaces as first proposed by Smolyak [88] and further investigated by Refs. [46, 8, 43, 105, 72, 73] that reduce dramatically the number of collocation points, while preserving a high level of accuracy.

Here we follow the notation and extend the description in Ref. [72] to describe the Smolyak *isotropic* formulas $\mathcal{A}(w, n)$, where w is a level that is independent of dimension³. The Smolyak formulas are just linear combinations of the product formulas in Eq. 2.38 with the following key property: only products with a relatively small number of points are used. With $\mathcal{U}^0 = 0$ and for $i \geq 1$ define

$$\Delta^i = \mathcal{U}^i - \mathcal{U}^{i-1}. \quad (2.39)$$

and we set $|\mathbf{i}| = i_1 + \dots + i_n$. Then the isotropic Smolyak quadrature formula is given by

$$\mathcal{A}(w, n) = \sum_{|\mathbf{i}| \leq w+n} (\Delta^{i_1} \otimes \dots \otimes \Delta^{i_n}). \quad (2.40)$$

This form is preferred for use in forming hierarchical interpolants as described in Sections 2.2.2 and 2.4.3. For nodal interpolants and polynomial chaos in sparse grids, the following equivalent form [97] is often more convenient since it collapses repeated index sets

$$\mathcal{A}(w, n) = \sum_{w+1 \leq |\mathbf{i}| \leq w+n} (-1)^{w+n-|\mathbf{i}|} \binom{n-1}{w+n-|\mathbf{i}|} \cdot (\mathcal{U}^{i_1} \otimes \dots \otimes \mathcal{U}^{i_n}). \quad (2.41)$$

For each index set $\mathbf{i}$ of levels, linear or nonlinear growth rules are used to define the corresponding one-dimensional quadrature orders. The following growth rules are employed for indices $i \geq 1$, where closed and open refer to the inclusion and exclusion of the bounds within an interval, respectively:

$$\text{closed nonlinear : } m = \begin{cases} 1 & i = 1 \\ 2^{i-1} + 1 & i > 1 \end{cases} \quad (2.42)$$

$$\text{open nonlinear : } m = 2^i - 1 \quad (2.43)$$

$$\text{open linear : } m = 2i - 1 \quad (2.44)$$

Nonlinear growth rules are used for fully nested rules (e.g., Clenshaw-Curtis is closed fully nested and Gauss-Patterson is open fully nested), and linear growth rules are best for standard Gauss rules that take advantage of, at most, “weak” nesting (e.g., reuse of the center point).

Examples of isotropic sparse grids, constructed from the fully nested Clenshaw-Curtis abscissas and the weakly-nested Gaussian abscissas are shown in Figure 2.1, where $\Omega = [-1, 1]^2$ and both Clenshaw-Curtis and Gauss-Legendre employ nonlinear growth⁴ from Eqs. 2.42 and 2.43, respectively. There, we consider a two-dimensional parameter space and a maximum level $w = 5$ (sparse grid $\mathcal{A}(5, 2)$). To see the reduction in function evaluations with respect to full tensor product grids, we also include a plot of the corresponding Clenshaw-Curtis isotropic full tensor grid having the same maximum number of points in each direction, namely $2^w + 1 = 33$.

³Other common formulations use a dimension-dependent level q where $q \geq n$. We use $w = q - n$, where $w \geq 0$ for all n .

⁴We prefer linear growth for Gauss-Legendre, but employ nonlinear growth here for purposes of comparison.

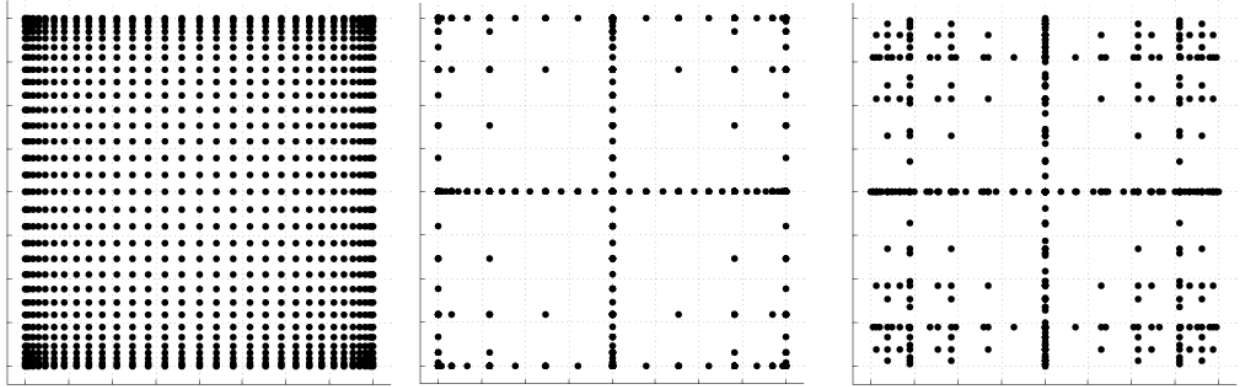


Figure 2.1: Two-dimensional grid comparison with a tensor product grid using Clenshaw-Curtis points (left) and sparse grids $\mathcal{A}(5, 2)$ utilizing Clenshaw-Curtis (middle) and Gauss-Legendre (right) points with nonlinear growth.

In [35], it is demonstrated that the synchronization of total-order PCE with the monomial resolution of a sparse grid is imperfect, and that sparse grid SC consistently outperforms sparse grid PCE when employing the sparse grid to directly evaluate the integrals in Eq. 2.35. In our Dakota implementation, we depart from the use of sparse integration of total-order expansions, and instead employ a linear combination of tensor expansions [21]. That is, we compute separate tensor polynomial chaos expansions for each of the underlying tensor quadrature grids (for which there is no synchronization issue) and then sum them using the Smolyak combinatorial coefficient (from Eq. 2.41 in the isotropic case). This improves accuracy, preserves the PCE/SC consistency property described in Section 2.6.2, and also simplifies PCE for the case of anisotropic sparse grids described next.

For anisotropic Smolyak sparse grids, a dimension preference vector is used to emphasize important stochastic dimensions. Given a mechanism for defining anisotropy, we can extend the definition of the sparse grid from that of Eq. 2.41 to weight the contributions of different index set components. First, the sparse grid index set constraint becomes

$$w\underline{\gamma} < \mathbf{i} \cdot \underline{\gamma} \leq w\underline{\gamma} + |\underline{\gamma}| \quad (2.45)$$

where $\underline{\gamma}$ is the minimum of the dimension weights γ_k , $k = 1$ to n . The dimension weighting vector $\underline{\gamma}$ amplifies the contribution of a particular dimension index within the constraint, and is therefore inversely related to the dimension preference (higher weighting produces lower index set levels). For the isotropic case of all $\gamma_k = 1$, it is evident that you reproduce the isotropic index constraint $w + 1 \leq |\mathbf{i}| \leq w + n$ (note the change from $<$ to $\leq$). Second, the combinatorial coefficient for adding the contribution from each of these index sets is modified as described in [15].

2.6.4 Cubature

Cubature rules [89, 104] are specifically optimized for multidimensional integration and are distinct from tensor-products and sparse grids in that they are not based on combinations of one-dimensional Gauss quadrature rules. They have the advantage of improved scalability to large numbers of random variables, but are restricted in integrand order and require homogeneous random variable sets (achieved via transformation). For example, optimal rules for integrands of 2, 3, and 5 and either Gaussian or uniform densities allow low-order polynomial chaos expansions ($p = 1$ or 2) that are useful for global sensitivity analysis including main effects and, for $p = 2$, all two-way interactions.

2.7 Linear regression

Regression based approaches attempt to find regularized or exact solutions of the linear system:

$$\Psi \alpha = R \quad (2.46)$$

to solve for the complete set of PCE coefficients α that best reproduce a set of response values R . The set of response values can be defined on an unstructured grid obtained from sampling within the density function of ξ (point collocation [96, 59]) or on a structured grid defined from uniform random sampling on the multi-index⁵ of a tensor-product quadrature grid (probabilistic collocation [91]), where the quadrature is of sufficient order to avoid sampling at roots of the basis polynomials⁶. In either case, each row of the matrix Ψ contains the N_t multivariate polynomial terms Ψ_j evaluated at a particular ξ sample. An over-sampling is most commonly used ([59] recommends $2N_t$ samples), resulting in a least squares solution for the over-determined system, although unique determination (N_t samples) and under-determination (fewer than N_t samples) are also supported. As for sampling-based coefficient estimation, this approach is only valid for PCE and does not require synchronization with monomial coverage; thus it is common to combine this coefficient estimation approach with a traditional total-order chaos expansion in order to keep sampling requirements low. In this case, simulation requirements for this approach scale as $\frac{r(n+p)!}{n!p!}$ (r is an over-sampling factor with typical values $0.1 \leq r \leq 2$), which can be significantly more affordable than isotropic tensor-product quadrature (scales as $(p+1)^n$ for standard Gauss rules) for larger problems. Additional regression equations can be obtained through the use of derivative information (gradients and Hessians) from each collocation point (refer to `use_derivatives` in the PCE regression specification details in the Dakota Reference Manual [2]), which can aid in scaling with respect to the number of random variables, particularly for adjoint-based derivative approaches. Finally, one can additionally modify the order of the exponent applied to N_t in the over-sampling ratio calculation (refer to `ratio_order` in the PCE regression specification details in the Dakota Reference Manual [2]).

Various methods can be employed to solve (2.46). The relative accuracy of each method is problem dependent. Traditionally, the most frequently used method has been least squares regression. However when Ψ is under-determined, minimizing the residual with respect to the ℓ_2 norm typically produces poor solutions. Compressed sensing methods have been successfully used to address this limitation [11, 29]. Such methods attempt to only identify the elements of the coefficient vector α with the largest magnitude and enforce as many elements as possible to be zero. Such solutions are often called sparse solutions. Dakota provides algorithms that solve the following formulations:

- Basis Pursuit (BP) [18]

$$\alpha = \arg \min \|\alpha\|_{\ell_1} \quad \text{such that} \quad \Psi \alpha = R \quad (2.47)$$

The BP solution is obtained in Dakota, by transforming (2.47) to a linear program which is then solved using the primal-dual interior-point method [13, 18].

- Basis Pursuit DeNoising (BPDN) [18].

$$\alpha = \arg \min \|\alpha\|_{\ell_1} \quad \text{such that} \quad \|\Psi \alpha - R\|_{\ell_2} \leq \varepsilon \quad (2.48)$$

The BPDN solution is computed in Dakota by transforming (2.48) to a quadratic cone problem which is solved using the log-barrier Newton method [13, 18].

When the matrix Ψ is not over-determined the BP and BPDN solvers used in Dakota will not return a solution. In such situations these methods simply return the least squares solution.

⁵Due to the discrete nature of index sampling, we enforce unique index samples by sorting and resampling as required.

⁶Generally speaking, dimension quadrature order m_i greater than dimension expansion order p_i .

- Orthogonal Matching Pursuit (OMP) [24],

$$\alpha = \arg \min \|\alpha\|_{\ell_0} \quad \text{such that} \quad \|\Psi\alpha - \mathbf{R}\|_{\ell_2} \leq \varepsilon \quad (2.49)$$

OMP is a heuristic method which greedily finds an approximation to (2.49). In contrast to the aforementioned techniques for solving BP and BPDN, which minimize an objective function, OMP constructs a sparse solution by iteratively building up an approximation of the solution vector α . The vector is approximated as a linear combination of a subset of active columns of Ψ . The active set of columns is built column by column, in a greedy fashion, such that at each iteration the inactive column with the highest correlation (inner product) with the current residual is added.

- Least Angle Regression (LARS) [32] and Least Absolute Shrinkage and Selection Operator (LASSO) [92]

$$\alpha = \arg \min \|\Psi\alpha - \mathbf{R}\|_{\ell_2}^2 \quad \text{such that} \|\alpha\|_{\ell_1} \leq \tau \quad (2.50)$$

A greedy solution can be found to (2.50) using the LARS algorithm. Alternatively, with only a small modification, one can provide a rigorous solution to this global optimization problem, which we refer to as the LASSO solution. Such an approach is identical to the homotopy algorithm of Osborne et al [76]. It is interesting to note that Efron [32] experimentally observed that the basic, faster LARS procedure is often identical to the LASSO solution.

The LARS algorithm is similar to OMP. LARS again maintains an active set of columns and again builds this set by adding the column with the largest correlation with the residual to the current residual. However, unlike OMP, LARS solves a penalized least squares problem at each step taking a step along an equiangular direction, that is, a direction having equal angles with the vectors in the active set. LARS and OMP do not allow a column (PCE basis) to leave the active set. However if this restriction is removed from LARS (it cannot be from OMP) the resulting algorithm can provably solve (2.50) and generates the LASSO solution.

- Elastic net [107]

$$\alpha = \arg \min \|\Psi\alpha - \mathbf{R}\|_{\ell_2}^2 \quad \text{such that} \quad (1 - \lambda)\|\alpha\|_{\ell_1} + \lambda\|\alpha\|_{\ell_2}^2 \leq \tau \quad (2.51)$$

The elastic net was developed to overcome some of the limitations of the LASSO formulation. Specifically: if the $(M \times N)$ Vandermonde matrix Ψ is over-determined ($M > N$), the LASSO selects at most N variables before it saturates, because of the nature of the convex optimization problem; if there is a group of variables among which the pairwise correlations are very high, then the LASSO tends to select only one variable from the group and does not care which one is selected; and finally if there are high correlations between predictors, it has been empirically observed that the prediction performance of the LASSO is dominated by ridge regression [92]. Here we note that it is hard to estimate the λ penalty in practice and the aforementioned issues typically do not arise very often when solving (2.46). The elastic net formulation can be solved with a minor modification of the LARS algorithm.

OMP and LARS add a PCE basis one step at a time. If α contains only k non-zero terms then these methods will only take k -steps. The homotopy version of LARS also adds only basis at each step, however it can also remove bases, and thus can take more than k steps. For some problems, the LARS and homotopy solutions will coincide. Each step of these algorithm provides a possible estimation of the PCE coefficients. However, without knowledge of the target function, there is no easy way to estimate which coefficient vector is best. With some additional computational effort (which will likely be minor to the cost of obtaining model simulations), cross validation can be used to choose an appropriate coefficient vector.

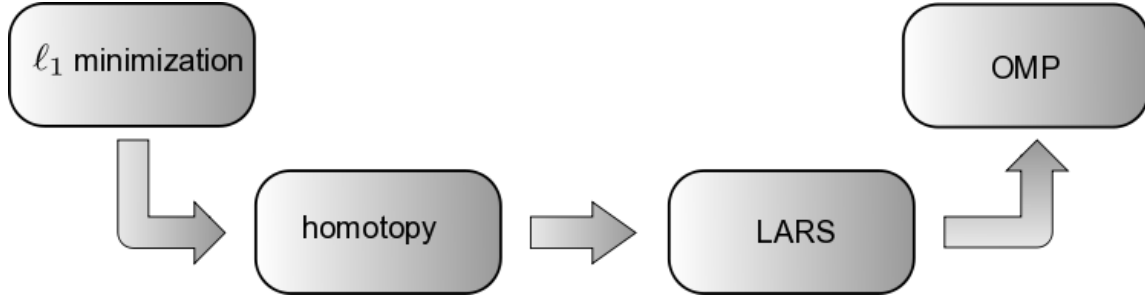


Figure 2.2: Bridging provably convergent ℓ_1 minimization algorithms and greedy algorithms such as OMP. (1) Homotopy provably solves ℓ_1 minimization problems [32]. (2) LARS is obtained from homotopy by removing the sign constraint check. (3) OMP and LARS are similar in structure, the only difference being that OMP solves a least-squares problem at each iteration, whereas LARS solves a linearly penalized least-squares problem. Figure and caption based upon Figure 1 in [28].

2.7.1 Cross validation

Cross validation can be used to find a coefficient vector α that approximately minimizes $\|\hat{f}(\mathbf{x}) - f(\mathbf{x})\|_{L^2(\rho)}$, where f is the target function and $\hat{f}$ is the PCE approximation using α . Given training data $\mathbf{X}$ and a set of algorithm parameters β (which can be step number in an algorithm such as OMP, or PCE maximum degree), K -folds cross validation divides $\mathbf{X}$ into K sets (folds) $\mathbf{X}_k$, $k = 1, \dots, K$ of equal size. A PCE $\hat{f}_\beta^{-k}(\mathbf{X})$, is built on the training data $\mathbf{X}_{\text{tr}} = \mathbf{X} \setminus \mathbf{X}_k$ with the k -th fold removed, using the tuning parameters β . The remaining data $\mathbf{X}_k$ is then used to estimate the prediction error. The prediction error is typically approximated by $e(\hat{f}) = \|\hat{f}(\mathbf{x}) - f(\mathbf{x})\|_{\ell_2}$, $\mathbf{x} \in \mathbf{X}_k$ [54]. This process is then repeated K times, removing a different fold from the training set each time.

The cross validation error is taken to be the average of the prediction errors for the K -experiments

$$CV(\hat{f}_\beta) = \mathbb{E}[e(\hat{f}_\beta^{-k})] = \frac{1}{K} \sum_{k=1}^K e(\hat{f}_\beta^{-k})$$

We minimize $CV(\hat{f}_\beta)$ as a surrogate for minimizing $\|\hat{f}_\beta(\mathbf{x}) - f(\mathbf{x})\|_{L^2(\rho)}$ and choose the tuning parameters

$$\beta^* = \arg \min CV(\hat{f}_\beta) \text{Var}[e(\hat{f}_\beta^{-k})] \quad (2.52)$$

to construct the final “best” PCE approximation of f that the training data can produce.

2.7.2 Orthogonal Least Interpolation

Orthogonal least interpolation (OLI) [70] enables the construction of interpolation polynomials based on arbitrarily located grids in arbitrary dimensions. The interpolation polynomials can be constructed using orthogonal polynomials corresponding to the probability distribution function of the uncertain random variables.

The algorithm for constructing an OLI is split into three stages: (i) basis determination - transform the orthogonal basis into a polynomial space that is amenable for interpolation; (ii) coefficient determination - determine the interpolatory coefficients on the transformed basis elements; (iii) connection problem - translate the coefficients on the transformed basis to coefficients in the original orthogonal basis. These three steps can be achieved by a sequence of LU and QR factorizations.

Orthogonal least interpolation is closely related to the aforementioned regression methods, in that OLI can be used to build approximations of simulation models when computing structured simulation data, such as sparse grids or cubature nodes, is infeasible. The interpolants produced by OLI have two additional important properties. Firstly, the orthogonal least interpolant is the lowest-degree polynomial that interpolates the data. Secondly, the least orthogonal interpolation space is monotonic. This second property means that the least interpolant can be extended to new data without the need to completely reconstructing the interpolant. The transformed interpolation basis can simply be extended to include the new necessary basis functions.

2.8 Analytic moments

Mean and covariance of polynomial chaos expansions are available in simple closed form:

$$\mu_i = \langle R_i \rangle \cong \sum_{k=0}^P \alpha_{ik} \langle \Psi_k(\xi) \rangle = \alpha_{i0} \quad (2.53)$$

$$\Sigma_{ij} = \langle (R_i - \mu_i)(R_j - \mu_j) \rangle \cong \sum_{k=1}^P \sum_{l=1}^P \alpha_{ik} \alpha_{jl} \langle \Psi_k(\xi) \Psi_l(\xi) \rangle = \sum_{k=1}^P \alpha_{ik} \alpha_{jk} \langle \Psi_k^2 \rangle \quad (2.54)$$

where the norm squared of each multivariate polynomial is computed from Eq. 2.36. These expressions provide exact moments of the expansions, which converge under refinement to moments of the true response functions.

Similar expressions can be derived for stochastic collocation:

$$\mu_i = \langle R_i \rangle \cong \sum_{k=1}^{N_p} r_{ik} \langle L_k(\xi) \rangle = \sum_{k=1}^{N_p} r_{ik} w_k \quad (2.55)$$

$$\Sigma_{ij} = \langle R_i R_j \rangle - \mu_i \mu_j \cong \sum_{k=1}^{N_p} \sum_{l=1}^{N_p} r_{ik} r_{jl} \langle L_k(\xi) L_l(\xi) \rangle - \mu_i \mu_j = \sum_{k=1}^{N_p} r_{ik} r_{jk} w_k - \mu_i \mu_j \quad (2.56)$$

where we have simplified the expectation of Lagrange polynomials constructed at Gauss points and then integrated at these same Gauss points. For tensor grids and sparse grids with fully nested rules, these expectations leave only the weight corresponding to the point for which the interpolation value is one, such that the final equalities in Eqs. 2.55–2.56 hold precisely. For sparse grids with non-nested rules, however, interpolation error exists at the collocation points, such that these final equalities hold only approximately. In this case, we have the choice of computing the moments based on sparse numerical integration or based on the moments of the (imperfect) sparse interpolant, where small differences may exist prior to numerical convergence. In Dakota, we employ the former approach; i.e., the right-most expressions in Eqs. 2.55–2.56 are employed for all tensor and sparse cases irregardless of nesting. Skewness and kurtosis calculations as well as sensitivity derivations in the following sections are also based on this choice. The expressions for skewness and (excess) kurtosis from direct numerical integration of the response function are as follows:

$$\gamma_{1i} = \left\langle \left(\frac{R_i - \mu_i}{\sigma_i} \right)^3 \right\rangle \cong \frac{1}{\sigma_i^3} \left[\sum_{k=1}^{N_p} (r_{ik} - \mu_i)^3 w_k \right] \quad (2.57)$$

$$\gamma_{2i} = \left\langle \left(\frac{R_i - \mu_i}{\sigma_i} \right)^4 \right\rangle - 3 \cong \frac{1}{\sigma_i^4} \left[\sum_{k=1}^{N_p} (r_{ik} - \mu_i)^4 w_k \right] - 3 \quad (2.58)$$

2.9 Local sensitivity analysis: derivatives with respect to expansion variables

Polynomial chaos expansions are easily differentiated with respect to the random variables [82]. First, using Eq. 2.20,

$$\frac{dR}{d\xi_i} = \sum_{j=0}^P \alpha_j \frac{d\Psi_j}{d\xi_i}(\xi) \quad (2.59)$$

and then using Eq. 2.19,

$$\frac{d\Psi_j}{d\xi_i}(\xi) = \frac{d\psi_{t_i^j}}{d\xi_i}(\xi_i) \prod_{\substack{k=1 \\ k \neq i}}^n \psi_{t_k^j}(\xi_k) \quad (2.60)$$

where the univariate polynomial derivatives $\frac{d\psi}{d\xi}$ have simple closed form expressions for each polynomial in the Askey scheme [1]. Finally, using the Jacobian of the (extended) Nataf variable transformation,

$$\frac{dR}{dx_i} = \frac{dR}{d\xi} \frac{d\xi}{dx_i} \quad (2.61)$$

which simplifies to $\frac{dR}{d\xi_i} \frac{d\xi_i}{dx_i}$ in the case of uncorrelated x_i .

Similar expressions may be derived for stochastic collocation, starting from Eq. 2.26:

$$\frac{dR}{d\xi_i} = \sum_{j=1}^{N_p} r_j \frac{dL_j}{d\xi_i}(\xi) \quad (2.62)$$

where the multidimensional interpolant L_j is formed over either tensor-product quadrature points or a Smolyak sparse grid. For the former case, the derivative of the multidimensional interpolant L_j involves differentiation of Eq. 2.27:

$$\frac{dL_j}{d\xi_i}(\xi) = \frac{dL_{c_i^j}}{d\xi_i}(\xi_i) \prod_{\substack{k=1 \\ k \neq i}}^n L_{c_k^j}(\xi_k) \quad (2.63)$$

and for the latter case, the derivative involves a linear combination of these product rules, as dictated by the Smolyak recursion shown in Eq. 2.41. Finally, calculation of $\frac{dR}{dx_i}$ involves the same Jacobian application shown in Eq. 2.61.

2.10 Global sensitivity analysis: variance-based decomposition

In addition to obtaining derivatives of stochastic expansions with respect to the random variables, it is possible to obtain variance-based sensitivity indices from the stochastic expansions. Variance-based sensitivity indices are explained in the Design of Experiments Chapter of the User's Manual [3]. The concepts are summarized here as well. Variance-based decomposition is a global sensitivity method that summarizes how the uncertainty in model output can be apportioned to uncertainty in individual input variables. VBD uses two primary measures, the main effect sensitivity index S_i and the total effect index T_i . These indices are also called the Sobol' indices. The main effect sensitivity index corresponds to the fraction of the uncertainty in the output, Y , that can be attributed to input x_i alone. The total effects index corresponds to the fraction of the uncertainty in the output, Y , that can be attributed to input x_i and its interactions with other variables. The main effect sensitivity index compares

the variance of the conditional expectation $Var_{x_i}[E(Y|x_i)]$ against the total variance $Var(Y)$. Formulas for the indices are:

$$S_i = \frac{Var_{x_i}[E(Y|x_i)]}{Var(Y)} \quad (2.64)$$

and

$$T_i = \frac{E(Var(Y|x_{-i}))}{Var(Y)} = \frac{Var(Y) - Var(E[Y|x_{-i}])}{Var(Y)} \quad (2.65)$$

where $Y = f(\mathbf{x})$ and $x_{-i} = (x_1, \dots, x_{i-1}, x_{i+1}, \dots, x_m)$.

The calculation of S_i and T_i requires the evaluation of m-dimensional integrals which are typically approximated by Monte-Carlo sampling. However, in stochastic expansion methods, it is possible to obtain the sensitivity indices as analytic functions of the coefficients in the stochastic expansion. The derivation of these results is presented in [90]. The sensitivity indices are printed as a default when running either polynomial chaos or stochastic collocation in Dakota. Note that in addition to the first-order main effects, S_i , we are able to calculate the sensitivity indices for higher order interactions such as the two-way interaction $S_{i,j}$.

2.11 Automated Refinement

Several approaches for refinement of stochastic expansions are currently supported, involving uniform or dimension-adaptive approaches to p- or h-refinement using structured (isotropic, anisotropic, or generalized) or unstructured grids. Specific combinations include:

- uniform refinement with unbiased grids
 - p-refinement: isotropic global tensor/sparse grids (PCE, SC) and regression (PCE only) with global basis polynomials
 - h-refinement: isotropic global tensor/sparse grids with local basis polynomials (SC only)
- dimension-adaptive refinement with biased grids
 - p-refinement: anisotropic global tensor/sparse grids with global basis polynomials using global sensitivity analysis (PCE, SC) or spectral decay rate estimation (PCE only)
 - h-refinement: anisotropic global tensor/sparse grids with local basis polynomials (SC only) using global sensitivity analysis
- goal-oriented dimension-adaptive refinement with greedy adaptation
 - p-refinement: generalized sparse grids with global basis polynomials (PCE, SC)
 - h-refinement: generalized sparse grids with local basis polynomials (SC only)

Each involves incrementing the global grid using differing grid refinement criteria, synchronizing the stochastic expansion specifics for the updated grid, and then updating the statistics and computing convergence criteria. Future releases will support local h-refinement approaches that can replace or augment the global grids currently supported. The sub-sections that follow enumerate each of the first level bullets above.

2.11.1 Uniform refinement with unbiased grids

Uniform refinement involves ramping the resolution of a global structured or unstructured grid in an unbiased manner and then refining an expansion to synchronize with this increased grid resolution. In the case of increasing the order of an isotropic tensor-product quadrature grid or the level of an isotropic Smolyak sparse grid, a p-refinement approach increases the order of the global basis polynomials (Sections 2.1, 2.2.1.1, and 2.2.1.2) in a synchronized manner and an h-refinement approach reduces the approximation range of fixed order local basis polynomials (Sections 2.2.1.3 and 2.2.1.4). And in the case of uniform p-refinement with PCE regression, the collocation oversampling ratio (refer to Methods specification within Dakota Reference Manual [2]) is held fixed, such that an increment in isotropic expansion order is matched with a corresponding increment in the number of structured (probabilistic collocation) or unstructured samples (point collocation) employed in the linear least squares solve.

For uniform refinement, anisotropic dimension preferences are not computed, and the only algorithmic requirements are:

- With the usage of nested integration rules with restricted exponential growth, Dakota must ensure that a change in level results in a sufficient change in the grid; otherwise premature convergence could occur within the refinement process. If no change is initially detected, Dakota continues incrementing the order/level (without grid evaluation) until the number of grid points increases.
- A convergence criterion is required. For uniform refinement, Dakota employs the L^2 norm of the change in the response covariance matrix as a general-purpose convergence metric.

2.11.2 Dimension-adaptive refinement with biased grids

Dimension-adaptive refinement involves ramping the order of a tensor-product quadrature grid or the level of a Smolyak sparse grid anisotropically, that is, using a defined dimension preference. This dimension preference may be computed from local sensitivity analysis, global sensitivity analysis, a posteriori error estimation, or decay rate estimation. In the current release, we focus on global sensitivity analysis and decay rate estimation. In the former case, dimension preference is defined from total Sobol' indices (Eq. 2.65) and is updated on every iteration of the adaptive refinement procedure, where higher variance attribution for a dimension indicates higher preference for that dimension. In the latter case, the spectral decay rates for the polynomial chaos coefficients are estimated using the available sets of univariate expansion terms (interaction terms are ignored). Given a set of scaled univariate coefficients (scaled to correspond to a normalized basis), the decay rate can be inferred using a technique analogous to Richardson extrapolation. The dimension preference is then defined from the inverse of the rate: slowly converging dimensions need greater refinement pressure. For both of these cases, the dimension preference vector supports anisotropic sparse grids based on a linear index-set constraint (Eq. 2.45) or anisotropic tensor grids (Eq. 2.38) with dimension order scaled proportionately to preference; for both grids, dimension refinement lower bound constraints are enforced to ensure that all previously evaluated points remain in new refined grids.

Given an anisotropic global grid, the expansion refinement proceeds as for the uniform case, in that the p-refinement approach increases the order of the global basis polynomials (Sections 2.1, 2.2.1.1, and 2.2.1.2) in a synchronized manner and an h-refinement approach reduces the approximation range of fixed order local basis polynomials (Sections 2.2.1.3 and 2.2.1.4). Also, the same grid change requirements and convergence criteria described for uniform refinement (Section 2.11.1) are applied in this case.

2.11.3 Goal-oriented dimension-adaptive refinement with greedy adaptation

Relative to the uniform and dimension-adaptive refinement capabilities described previously, the generalized sparse grid algorithm [47] supports greater flexibility in the definition of sparse grid index sets and supports refinement controls based on general statistical quantities of interest (QOI). This algorithm was originally intended for adaptive numerical integration on a hypercube, but can be readily extended to the adaptive refinement of stochastic expansions using the following customizations:

- In addition to hierarchical interpolants in SC, we employ independent polynomial chaos expansions for each active and accepted index set. Pushing and popping index sets then involves increments of tensor chaos expansions (as described in Section 2.6.3) along with corresponding increments to the Smolyak combinatorial coefficients.
- Since we support bases for more than uniform distributions on a hypercube, we exploit rule nesting when possible (i.e., Gauss-Patterson for uniform or transformed uniform variables, and Genz-Keister for normal or transformed normal variables), but we do not require it. This implies a loss of some algorithmic simplifications in [47] that occur when grids are strictly hierarchical.
- In the evaluation of the effect of a trial index set, the goal in [47] is numerical integration and the metric is the size of the increment induced by the trial set on the expectation of the function of interest. It is straightforward to instead measure the effect of a trial index set on response covariance, numerical probability, or other statistical QOI computed by post-processing the resulting PCE or SC expansion. By tying the refinement process more closely to the statistical QOI, the refinement process can become more efficient in achieving the desired analysis objectives.

Hierarchical increments in a variety of statistical QoI may be derived, starting from increments in response mean and covariance. The former is defined from computing the expectation of the difference interpolants in Eqs. 2.31-2.32, and the latter is defined as:

$$\Delta\Sigma_{ij} = \Delta E[R_i R_j] - \mu_{R_i} \Delta E[R_j] - \mu_{R_j} \Delta E[R_i] - \Delta E[R_i] \Delta E[R_j] \quad (2.66)$$

Increments in standard deviation and reliability indices can subsequently be defined, where care is taken to preserve numerical precision through the square root operation (e.g., via Boost sqrt1pml()).

Given these customizations, the algorithmic steps can be summarized as:

1. *Initialization:* Starting from an initial isotropic or anisotropic reference grid (often the $w = 0$ grid corresponding to a single collocation point), accept the reference index sets as the old set and define active index sets using the admissible forward neighbors of all old index sets.
2. *Trial set evaluation:* Evaluate the tensor grid corresponding to each trial active index set, form the tensor polynomial chaos expansion or tensor interpolant corresponding to it, update the Smolyak combinatorial coefficients, and combine the trial expansion with the reference expansion. Perform necessary bookkeeping to allow efficient restoration of previously evaluated tensor expansions.
3. *Trial set selection:* Select the trial index set that induces the largest change in the statistical QOI, normalized by the cost of evaluating the trial index set (as indicated by the number of new collocation points in the trial grid). In our implementation, the statistical QOI is defined using an L^2 norm of change in CDF/CCDF probability/reliability/response level mappings, when level mappings are present, or L^2 norm of change in response covariance, when level mappings are not present.

4. *Update sets:* If the largest change induced by the trial sets exceeds a specified convergence tolerance, then promote the selected trial set from the active set to the old set and update the active sets with new admissible forward neighbors; return to step 2 and evaluate all trial sets with respect to the new reference point. If the convergence tolerance is satisfied, advance to step 5.
5. *Finalization:* Promote all remaining active sets to the old set, update the Smolyak combinatorial coefficients, and perform a final combination of tensor expansions to arrive at the final result for the statistical QOI.

2.12 Multifidelity methods

In a multifidelity uncertainty quantification approach employing stochastic expansions, we seek to utilize a predictive low-fidelity model to our advantage in reducing the number of high-fidelity model evaluations required to compute high-fidelity statistics to a particular precision. When a low-fidelity model captures useful trends of the high-fidelity model, then the model discrepancy may have either lower complexity, lower variance, or both, requiring less computational effort to resolve its functional form than that required for the original high-fidelity model [71].

To accomplish this goal, an expansion will first be formed for the model discrepancy (the difference between response results if additive correction or the ratio of results if multiplicative correction). These discrepancy functions are the same functions approximated in surrogate-based minimization (see Surrogate Models section within the Models chapter of the User's Manual [3]). The exact discrepancy functions are

$$A(\xi) = R_{hi}(\xi) - R_{lo}(\xi) \quad (2.67)$$

$$B(\xi) = \frac{R_{hi}(\xi)}{R_{lo}(\xi)} \quad (2.68)$$

Approximating the high-fidelity response functions using approximations of these discrepancy functions then involves

$$\hat{R}_{hi_A}(\xi) = R_{lo}(\xi) + \hat{A}(\xi) \quad (2.69)$$

$$\hat{R}_{hi_B}(\xi) = R_{lo}(\xi) \hat{B}(\xi) \quad (2.70)$$

where $\hat{A}(\xi)$ and $\hat{B}(\xi)$ are stochastic expansion approximations to the exact correction functions:

$$\hat{A}(\xi) = \sum_{j=0}^{P_{hi}} \alpha_j \Psi_j(\xi) \text{ or } \sum_{j=1}^{N_{hi}} a_j L_j(\xi) \quad (2.71)$$

$$\hat{B}(\xi) = \sum_{j=0}^{P_{hi}} \beta_j \Psi_j(\xi) \text{ or } \sum_{j=1}^{N_{hi}} b_j L_j(\xi) \quad (2.72)$$

where α_j and β_j are the spectral coefficients for a polynomial chaos expansion (evaluated via Eq. 2.35, for example) and a_j and b_j are the interpolation coefficients for stochastic collocation (values of the exact discrepancy evaluated at the collocation points).

Second, an expansion will be formed for the low fidelity surrogate model, where the intent is for the level of resolution to be higher than that required to resolve the discrepancy ($P_{lo} \gg P_{hi}$ or $N_{lo} \gg N_{hi}$; either enforced statically through order/level selections or automatically through adaptive refinement):

$$R_{lo}(\xi) \cong \sum_{j=0}^{P_{lo}} \gamma_j \Psi_j(\xi) \text{ or } \sum_{j=1}^{N_{lo}} r_{lo_j} L_j(\xi) \quad (2.73)$$

Then the two expansions are combined (added or multiplied) into a new expansion that approximates the high fidelity model, from which the final set of statistics are generated. For polynomial chaos expansions, this combination entails:

- in the additive case, the high-fidelity expansion is formed by simply overlaying the expansion forms and adding the spectral coefficients that correspond to the same basis polynomials.
- in the multiplicative case, the form of the high-fidelity expansion must first be defined to include all polynomial orders indicated by the products of each of the basis polynomials in the low fidelity and discrepancy expansions (most easily estimated from total-order, tensor, or sum of tensor expansions which involve simple order additions). Then the coefficients of this product expansion are computed as follows (shown generically for $z = xy$ where x , y , and z are each expansions of arbitrary form):

$$\sum_{k=0}^{P_z} z_k \Psi_k(\xi) = \sum_{i=0}^{P_x} \sum_{j=0}^{P_y} x_i y_j \Psi_i(\xi) \Psi_j(\xi) \quad (2.74)$$

$$z_k = \frac{\sum_{i=0}^{P_x} \sum_{j=0}^{P_y} x_i y_j \langle \Psi_i \Psi_j \Psi_k \rangle}{\langle \Psi_k^2 \rangle} \quad (2.75)$$

where tensors of one-dimensional basis triple products $\langle \psi_i \psi_j \psi_k \rangle$ are typically sparse and can be efficiently precomputed using one dimensional quadrature for fast lookup within the multidimensional triple products.

For stochastic collocation, the high-fidelity expansion generated from combining the low fidelity and discrepancy expansions retains the polynomial form of the low fidelity expansion, for which only the coefficients are updated in order to interpolate the sum or product values (and potentially their derivatives). Since we will typically need to produce values of a less resolved discrepancy expansion on a more resolved low fidelity grid to perform this combination, we utilize the discrepancy expansion rather than the original discrepancy function values for both interpolated and non-interpolated point values (and derivatives), in order to ensure consistency.

Chapter 3

Epistemic Methods

This chapter covers theoretical aspects of methods for propagating epistemic uncertainty.

3.1 Dempster-Shafer theory of evidence (DSTE)

In Dempster-Shafer theory, the event space is defined by a triple $(\mathcal{S}, \mathbb{S}, m)$ which defines $\mathcal{S}$ the universal set, $\mathbb{S}$ a countable collection of subsets of $\mathcal{S}$, and a notional measure m . $\mathcal{S}$ and $\mathbb{S}$ have a similar meaning to that in classical probability theory; the main difference is that $\mathbb{S}$, also known as the focal elements, does not have to be a σ -algebra over $\mathcal{S}$. The operator m is defined to be

$$m(\mathcal{U}) = \begin{cases} > 0 & \text{if } \mathcal{U} \in \mathbb{S} \\ 0 & \text{if } \mathcal{U} \subset \mathcal{S} \text{ and } \mathcal{U} \notin \mathbb{S} \end{cases} \quad (3.1)$$

$$\sum_{\mathcal{U} \in \mathbb{S}} m(\mathcal{U}) = 1 \quad (3.2)$$

where $m(\mathcal{U})$ is known as the basic probability assignment (BPA) of the set $\mathcal{U}$. In the DSTE framework, belief and plausibility are defined as:

$$\text{Bel}(\mathcal{E}) = \sum_{\{\mathcal{U} \mid \mathcal{U} \subset \mathcal{E}, \mathcal{U} \in \mathbb{S}\}} m(\mathcal{U}) \quad (3.3)$$

$$\text{Pl}(\mathcal{E}) = \sum_{\{\mathcal{U} \mid \mathcal{U} \cap \mathcal{E} \neq \emptyset, \mathcal{U} \in \mathbb{S}\}} m(\mathcal{U}) \quad (3.4)$$

The belief $\text{Bel}(\mathcal{E})$ is interpreted to be the minimum likelihood that is associated with the event $\mathcal{E}$. Similarly, the plausibility $\text{Pl}(\mathcal{E})$ is the maximum amount of likelihood that could be associated with $\mathcal{E}$. This particular structure allows us to handle unconventional inputs, such as conflicting pieces of evidence (e.g. dissenting expert opinions), that would be otherwise discarded in an interval analysis or probabilistic framework. The ability to make use of this information results in a commensurately more informed output.

The procedure to compute belief structures involves four major steps:

1. Determine the set of d -dimensional hypercubes that have a nonzero evidential measure
2. Compute the composite evidential measure (BPA) of each hypercube
3. Propagate each hypercube through the model and obtain the response bounds within each hypercube
4. Aggregate the minimum and maximum values of the response per hypercube with the BPAs to obtain cumulative belief and plausibility functions on the response (e.g. calculate a belief structure on the response).

The first step involves identifying combinations of focal elements defined on the inputs that define a hypercube. The second step involves defining an aggregate BPA for that hypercube, which is the product of the BPAs of the individual focal elements defining the hypercube. The third step involves finding the maximum and minimum values of the response value in each hypercube, and this part can be very computationally expensive. Finally, the results over all hypercubes are aggregated to form belief structures on the response.

Chapter 4

Surrogate Models

This chapter deals with the theory behind Dakota's surrogate models, which are also known as response surfaces and meta-models.

4.1 Kriging and Gaussian Process Models

In this discussion of Kriging and Gaussian Process (GP) models, vectors are indicated by a single underline and matrices are indicated by a double underline. Capital letters indicate data, or functions of data, that is used to construct an emulator. Lower case letters indicate arbitrary points, i.e. points where the simulator may or may not have been evaluated, and functions of arbitrary points. Estimates, approximations, and models are indicated by hats. For instance, $\hat{f}(\underline{x})$ is a model/emulator of the function $f(\underline{x})$ and $\hat{y}$ is the emulator's prediction or estimate of the true response $y = f(\underline{x})$ evaluated at the point $\underline{x}$. A tilde indicates a reordering of points/equations, with the possible omission of some but not all points. N is the number of points in the sample design and M is the number of input dimensions.

4.1.1 Kriging & Gaussian Processes: Function Values Only

The set of interpolation techniques known as Kriging, also referred to as Gaussian Processes, were originally developed in the geostatistics and spatial statistics communities to produce maps of underground geologic deposits based on a set of widely and irregularly spaced borehole sites[23]. Building a Kriging model typically involves the

1. Choice of a trend function,
2. Choice of a correlation function, and
3. Estimation of correlation parameters.

A Kriging emulator, $\hat{f}(\underline{x})$, consists of a trend function (frequently a least squares fit to the data, $\underline{g}(\underline{x})^T \underline{\beta}$) plus a Gaussian process error model, $\epsilon(\underline{x})$, that is used to correct the trend function.

$$\hat{f}(\underline{x}) = \underline{g}(\underline{x})^T \underline{\beta} + \epsilon(\underline{x})$$

This represents an estimated distribution for the unknown true surface, $f(\underline{x})$. The error model, $\epsilon(\underline{x})$, makes an adjustment to the trend function so that the emulator will interpolate, and have zero uncertainty at, the data points it was built from. The covariance between the error at two arbitrary points, $\underline{x}$ and $\underline{x}'$, is modeled as

$$\text{Cov}(y(\underline{x}), y(\underline{x}')) = \text{Cov}(\epsilon(\underline{x}), \epsilon(\underline{x}')) = \sigma^2 r(\underline{x}, \underline{x}').$$

Here σ^2 is known as the unadjusted variance and $r(\underline{x}, \underline{x}')$ is a correlation function. Measurement error can be modeled explicitly by modifying this to

$$\text{Cov}(\epsilon(\underline{x}), \epsilon(\underline{x}')) = \sigma^2 r(\underline{x}, \underline{x}') + \Delta^2 \delta(\underline{x} - \underline{x}')$$

where

$$\delta(\underline{x} - \underline{x}') = \begin{cases} 1 & \text{if } \underline{x} - \underline{x}' = \underline{0} \\ 0 & \text{otherwise} \end{cases}$$

and Δ^2 is the variance of the measurement error. In this work, the term “nugget” refers to the ratio $\eta = \frac{\Delta^2}{\sigma^2}$.

By convention, the terms simple Kriging, ordinary Kriging, and universal Kriging are used to indicate the three most common choices for the trend function. In simple Kriging, the trend is treated as a known constant, usually zero, $g(\underline{x})\beta \equiv 0$. Universal Kriging [68] uses a general polynomial trend model $g(\underline{x})^T \underline{\beta}$ whose coefficients are determined by least squares regression. Ordinary Kriging is essentially universal Kriging with a trend order of zero, i.e. the trend function is treated as an unknown constant and $g(\underline{x}) = 1$. N_β is the number of basis functions in $g(\underline{x})$ and therefore number of elements in the vector $\underline{\beta}$.

For a finite number of sample points, N , there will be uncertainty about the most appropriate value of the vector, $\underline{\beta}$, which can therefore be described as having a distribution of possible values. If one assumes zero prior knowledge about this distribution, which is referred to as the “vague prior” assumption, then the maximum likelihood value of $\underline{\beta}$ can be computed via least squares generalized by the inverse of the error model’s correlation matrix, $\underline{R}$

$$\hat{\underline{\beta}} = (\underline{G}^T \underline{R}^{-1} \underline{G})^{-1} (\underline{G}^T \underline{R}^{-1} \underline{Y}).$$

Here $\underline{G}$ is a N by N_β matrix that contains the evaluation of the least squares basis functions at all points in $\underline{X}$, such that $G_{i,l} = g_l(\underline{X}_i)$. The real, symmetric, positive-definite correlation matrix, $\underline{R}$, of the error model contains evaluations of the correlation function, r , at all pairwise combination of points (rows) in the sample design, $\underline{X}$.

$$R_{i,j} = R_{j,i} = r(\underline{X}_i, \underline{X}_j) = r(\underline{X}_j, \underline{X}_i)$$

There are many options for r , among them are the following families of correlation functions:

- **Powered-Exponential**

$$r(\underline{X}_i, \underline{X}_j) = \exp\left(-\sum_{k=1}^M \theta_k |X_{i,k} - X_{j,k}|^\gamma\right) \quad (4.1)$$

where $0 < \gamma \leq 2$ and $0 < \theta_k$.

- **Matern**

$$r(\underline{X}_i, \underline{X}_j) = \prod_{k=1}^M \frac{2^{1-\nu}}{\Gamma(\nu)} (\theta_k |X_{i,k} - X_{j,k}|)^\nu \mathcal{K}_\nu(\theta_k |X_{i,k} - X_{j,k}|)$$

where $0 < \nu$, $0 < \theta_k$, and $\mathcal{K}_\nu(\cdot)$ is the modified Bessel function of order ν ; $\nu = s + \frac{1}{2}$ is the smallest value of ν which results in a Kriging model that is s times differentiable.

- **Cauchy**

$$r(\underline{X}_i, \underline{X}_j) = \prod_{k=1}^M (1 + \theta_k |X_{i,k} - X_{j,k}|^\gamma)^{-\nu}$$

where $0 < \gamma \leq 2$, $0 < \nu$, and $0 < \theta_k$.

Gneiting et al. [51] provide a more thorough discussion of the properties of and relationships between these three families. Some additional correlation functions include the Dagum family [9] and cubic splines.

The squared exponential or Gaussian correlation function (Equation 4.1 with $\gamma = 2$) was selected to be the first correlation function implemented in Dakota on the basis that its infinite smoothness or differentiability should aid in leveraging the anticipated and hopefully sparse data. For the Gaussian correlation function, the correlation parameters, $\underline{\theta}$, are related to the correlation lengths, $\underline{L}$, by

$$\theta_k = \frac{1}{2 L_k^2}. \quad (4.2)$$

Here, the correlation lengths, $\underline{L}$, are analogous to standard deviations in the Gaussian or normal distribution and often have physical meaning. The adjusted (by data) mean of the emulator is a best linear unbiased estimator of the unknown true function,

$$\hat{y} = E(\hat{f}(\underline{x}) | \underline{f}(\underline{X})) = \underline{g}(\underline{x})^T \hat{\underline{\beta}} + \underline{r}(\underline{x})^T \underline{R}^{-1} \underline{\epsilon}. \quad (4.3)$$

Here, $\underline{\epsilon} = (\underline{Y} - \underline{G} \hat{\underline{\beta}})$ is the known vector of differences between the true outputs and trend function at all points in $\underline{X}$ and the vector $\underline{r}(\underline{x})$ is defined such that $r_i(\underline{x}) = r(\underline{x}, X_i)$. This correction can be interpreted as the projection of prior belief (the least squares fit) into the span of the data. The adjusted mean of the emulator will interpolate the data that the Kriging model was built from as long as its correlation matrix, $\underline{R}$, is numerically non-singular.

Ill-conditioning of $\underline{R}$ and other matrices is a recognized as a significant challenge for Kriging. Davis and Morris [25] gave a thorough review of six factors affecting the condition number of matrices associated with Kriging (from the perspective of semivariograms rather than correlation functions). They concluded that “Perhaps the best advice we can give is to be mindful of the condition number when building and solving kriging systems.”

In the context of estimating the optimal $\underline{\theta}$, Martin [67] stated that Kriging’s “three most prevalent issues are (1) ill-conditioned correlation matrices, (2) multiple local optimum, and (3) long ridges of near optimal values.” Because of the second issue, global optimization methods are more robust than local methods. Martin used constrained optimization to address ill-conditioning of $\underline{R}$.

Rennen [83] advocated that ill-conditioning be handled by building Kriging models from a uniform subset of available sample points. That option has been available in Dakota’s “Gaussian process” model (a separate implementation from Dakota’s “Kriging” model) since version 4.1 [40]. Note that Kriging/Gaussian-Process models will not exactly interpolate the discarded points. The implicit justification for this type of approach is that the row or columns of an ill-conditioned matrix contain a significant amount of duplicate information, and that when discarded, duplicate information should be easy to predict.

As of version 5.2, Dakota’s `kriging` model has a similar “discard near duplicate points” capability. However, it explicitly addresses the issue of unique information content. Points are **not** discarded prior to the construction

of the Kriging model. Instead, for each vector $\underline{\theta}$ examined that results in an ill-conditioned correlation matrix, $\underline{R}$, a pivoted Cholesky factorization of $\underline{R}$ is performed. This ranks the points according to how much unique information they contain. Note that the definition of information content depends on $\underline{\theta}$. Low information points are then discarded until $\underline{R}$ is no longer ill-conditioned, i.e. until it tightly meets a constraint on condition number. This can be done efficiently using a bisection search that calls LAPACK's fast estimate of the (reciprocal of the) condition number. The possibly, and often, improper subset of points used to construct the Kriging model is the one associated with the chosen $\underline{\theta}$. Methods for selecting $\underline{\theta}$ are discussed below. Since the points that are discarded are the ones that contain the least unique information, they are the ones that are easiest to predict and provide maximum improvement to the condition number.

Adding a nugget, η , to the diagonal entries of $\underline{R}$ is a popular approach for both accounting for measurement error in the data and alleviating ill-conditioning. However, doing so will cause the Kriging model to smooth or approximate rather than interpolate the data. Methods for choosing a nugget include:

- Choosing a nugget based on the variance of measurement error (if any); this will be an iterative process if σ^2 is not known in advance.
- Iteratively adding a successively larger nugget until $\underline{R} + \eta \underline{I}$ is no longer ill-conditioned.
- Exactly calculating the minimum nugget needed for a target condition number from $\underline{R}$'s maximum λ_{max} and minimum λ_{min} eigenvalues. The condition number of $\underline{R} + \eta \underline{I}$ is $\frac{\lambda_{max} + \eta}{\lambda_{min} + \eta}$. However, calculating eigenvalues is computationally expensive. Since Kriging's $\underline{R}$ matrix has all ones on the diagonal, its trace and therefore sum of eigenvalues is N . Consequently, a nugget value of $\eta = \frac{N}{\text{target condition number} - 1}$ will always alleviate ill-conditioning. A smaller nugget that is also guaranteed to alleviate ill-conditioning can be calculated from LAPACK's fast estimate of the reciprocal of $\underline{R}$'s condition number, $\text{rcond}(\underline{R})$.
- Treating η as another parameter to be selected by the same process used to choose $\underline{\theta}$. Two such approaches are discussed below.

The Kriging model's adjusted variance is commonly used as a spatially varying measure of uncertainty. Knowing where, and by how much, the model "doubts" its own predictions helps build user confidence in the predictions and can be utilized to guide the selection of new sample points during optimization or to otherwise improve the surrogate. The adjusted variance is

$$\begin{aligned} \text{Var}(\hat{y}) &= \text{Var}\left(\hat{f}(\underline{x}) | f(\underline{X})\right) \\ &= \hat{\sigma}^2 \left(1 - \underline{r}(\underline{x})^T \underline{R}^{-1} \underline{r}(\underline{x}) + \dots \right. \\ &\quad \left. \left(\underline{g}(\underline{x})^T - \underline{r}(\underline{x})^T \underline{R}^{-1} \underline{G}\right) \left(\underline{G}^T \underline{R}^{-1} \underline{G}\right)^{-1} \left(\underline{g}(\underline{x})^T - \underline{r}(\underline{x})^T \underline{R}^{-1} \underline{G}\right)^T\right) \end{aligned}$$

where the maximum likelihood estimate of the unadjusted variance is

$$\hat{\sigma}^2 = \frac{\underline{\epsilon}^T \underline{R}^{-1} \underline{\epsilon}}{N - N_\beta}.$$

There are two types of numerical approaches to choosing $\underline{\theta}$. One of these is to use Bayesian techniques such as Markov Chain Monte Carlo to obtain a distribution represented by an ensemble of vectors $\underline{\theta}$. In this case, evaluating the emulator's mean involves taking a weighted average of Equation 4.3 over the ensemble of $\underline{\theta}$ vectors.

The other, more common, approach to constructing a Kriging model involves using optimization to find the set of correlation parameters $\underline{\theta}$ that maximizes the likelihood of the model given the data. Dakota's `gaussian_process` and `kriging` models use the maximum likelihood approach. It is equivalent, and more convenient to maximize the natural logarithm of the likelihood, which assuming a vague prior is,

$$\log(\text{lik}(\underline{\theta})) = -\frac{1}{2} \left((N - N_\beta) \left(\frac{\hat{\sigma}^2}{\sigma^2} + \log(\sigma^2) + \log(2\pi) \right) + \dots \right. \\ \left. \log(\det(\underline{R})) + \log(\det(\underline{G}^T \underline{R}^{-1} \underline{G})) \right).$$

And, if one substitutes the maximum likelihood estimate $\hat{\sigma}^2$ in for σ^2 , then it is equivalent to minimize the following objective function

$$\text{obj}(\underline{\theta}) = \log(\hat{\sigma}^2) + \frac{\log(\det(\underline{R})) + \log(\det(\underline{G}^T \underline{R}^{-1} \underline{G}))}{N - N_\beta}.$$

Because of the division by $N - N_\beta$, this “per-equation” objective function is mostly independent of the number of sample points, N . It is therefore useful for comparing the (estimated) “goodness” of Kriging models that have different numbers of sample points, e.g. when an arbitrary number of points can be discarded by the pivoted Cholesky approach described above.

Note that the determinant of $\underline{R}$ (and $(\underline{G}^T \underline{R}^{-1} \underline{G})$) can be efficiently calculated as the square of the product of the diagonals of its Cholesky factorization. However, this will often underflow, i.e. go to zero, making its log incorrectly go to $-\infty$. A more accurate and robust calculation of $\log(\det(\underline{R}))$ can be achieved by taking twice the sum of the log of the diagonals of $\underline{R}$'s Cholesky factorization.

Also note, that in the absence of constraints, maximizing the likelihood would result in singular $\underline{R}$ which makes the emulator incapable of reproducing the data from which it was built. This is because a singular $\underline{R}$ makes $\log(\det(\underline{R})) = -\infty$ and the *estimate* of likelihood infinite. Constraints are therefore required. Two types of constraints are used in Dakota's `kriging` models.

The first of these is an explicit constraint on LAPACK's fast estimate of the (reciprocal of the) condition number, $2^{-40} < \text{rcond}(\underline{R})$. The value 2^{-40} was chosen based on the assumption that double precision arithmetic is used. Double precision numbers have 52 bits of precision. Therefore the $2^{-40} < \text{rcond}(\det(\underline{R}))$ implies that at least the leading three significant figures should be uncorrupted by round off error. In Dakota 5.2, this constraint is used to determine how many points can be retained in the pivoted Cholesky approach to subset selection described above.

The second, is a box constraint defining a small “feasible” region in correlation length space to search during the maximum likelihood optimization. Many global optimizers, such as the DIRECT (DIvision of RECTangles) used by Dakota's Gaussian Process (as the only option) and Kriging (as the default option) models, require a box constraint definition for the range of acceptable parameters. By default, Dakota's `kriging` model defines the input space to be the smallest hyper-rectangle that contains the sample design. The user has the option to define a larger input space that includes a region where they wish to extrapolate. Note that the emulator can be evaluated at points outside the defined input space, but this definition helps to determine the extent of the “feasible” region of correlation lengths. Let the input space be normalized to the unit hypercube centered at the origin. The average

distance between nearest neighboring points is then

$$d = \left(\frac{1}{N} \right)^{1/M}.$$

Dakota's "feasible" range of correlation lengths, $\underline{L}$, for the Gaussian correlation function is

$$\frac{d}{4} \leq L_k \leq 8d.$$

This range was chosen based on correlation lengths being analogous to the standard deviation in the Gaussian or Normal distribution. If the correlation lengths are set to $L_k = d/4$, then nearest neighboring points "should be" roughly four "standard deviations" away making them almost completely uncorrelated with each other. $\underline{R}$ would then be a good approximation of the identity matrix and have a condition number close to one. In the absence of a pathological spacing of points, this range of $\underline{L}$ should contain some non-singular $\underline{R}$. $L_k = 8d$ implies approximately 32% trust in what points 8 neighbors away have to say and 5% trust in what points 16 neighbors away have to say. It is possible that the optimal correlation lengths are larger than $8d$; but if so, then either almost all of the same information will be contained in more nearby neighbors, or it was not appropriate to use the squared-exponential/Gaussian correlation function. When other correlation functions are added to the Dakota Kriging implementation, each will be associated with its own range of appropriate correlation lengths chosen by similar reasoning. A different definition of d could be used for non-hypercube input domains.

4.1.2 Gradient Enhanced Kriging

This section focuses on the incorporation of derivative information into Kriging models and challenges in their implementation. Special attention is paid to conditioning issues.

There are at least three basic approaches for incorporating derivative information into Kriging. These are

1. **Indirect:** The sample design is augmented with fictitious points nearby actual sample points which are predicted from derivative information and then a Kriging model is built from the augmented design.
2. **Co-Kriging:** The derivatives with respect to each input variables are treated as separate but correlated output variables and a Co-Kriging model is built for the set of output variables. This would use $\binom{M+2}{2}$ $\underline{\theta}$ vectors.
3. **Direct:** The relationship between the response value and its derivatives is leveraged to use a single $\underline{\theta}$ by assuming

$$\text{Cov} \left(y(\underline{x}^1), \frac{\partial y(\underline{x}^2)}{\partial x_k^2} \right) = \frac{\partial}{\partial x_k^2} (\text{Cov}(y(\underline{x}^1), y(\underline{x}^2))). \quad (4.4)$$

Dakota 5.2 and later includes an implementation of the direct approach, herein referred to simply as Gradient Enhanced (universal) Kriging (GEK). The equations for GEK can be derived by assuming Equation 4.4 and then taking the same steps used to derive function value only Kriging. The superscript on $\underline{x}$ in Equation 4.4 and below indicates whether it's the 1st or 2nd input to $r(\underline{x}^1, \underline{x}^2)$. Note that when the first and second arguments are the same, the derivative of $r(\cdot, \cdot)$ with respect to the first argument is equal in magnitude but opposite in sign compared to the derivative with respect to the second argument. The GEK equations can also be obtained by

starting from a Kriging model and making the following substitutions $\underline{Y} \rightarrow \underline{Y}_{\nabla}$, $\underline{G} \rightarrow \underline{G}_{\nabla}$, $\underline{r} \rightarrow \underline{r}_{\nabla}$, $\underline{R} \rightarrow \underline{R}_{\nabla}$, and $N \rightarrow N_{\nabla} = N(1 + M)$, where N_{∇} is the number of equations rather than the number of points,

$$\underline{Y}_{\nabla} = \begin{bmatrix} \underline{Y} \\ \frac{\partial \underline{Y}}{\partial X_{:,1}} \\ \frac{\partial \underline{Y}}{\partial X_{:,2}} \\ \vdots \\ \frac{\partial \underline{Y}}{\partial X_{:,M}} \end{bmatrix}, \quad \underline{G}_{\nabla} = \begin{bmatrix} \underline{G} \\ \frac{\partial \underline{G}}{\partial X_{:,1}} \\ \frac{\partial \underline{G}}{\partial X_{:,2}} \\ \vdots \\ \frac{\partial \underline{G}}{\partial X_{:,M}} \end{bmatrix}, \quad \underline{r}_{\nabla} = \begin{bmatrix} \underline{r} \\ \frac{\partial \underline{r}}{\partial X_{:,1}} \\ \frac{\partial \underline{r}}{\partial X_{:,2}} \\ \vdots \\ \frac{\partial \underline{r}}{\partial X_{:,M}} \end{bmatrix}$$

$$\underline{R}_{\nabla} = \begin{bmatrix} \underline{R} & \frac{\partial \underline{R}}{\partial X_{:,1}^2} & \frac{\partial \underline{R}}{\partial X_{:,2}^2} & \cdots & \frac{\partial \underline{R}}{\partial X_{:,M}^2} \\ \frac{\partial \underline{R}}{\partial X_{:,1}^1} & \frac{\partial^2 \underline{R}}{\partial X_{:,1}^1 \partial X_{:,1}^2} & \frac{\partial^2 \underline{R}}{\partial X_{:,1}^1 \partial X_{:,2}^2} & \cdots & \frac{\partial^2 \underline{R}}{\partial X_{:,1}^1 \partial X_{:,M}^2} \\ \frac{\partial \underline{R}}{\partial X_{:,2}^1} & \frac{\partial^2 \underline{R}}{\partial X_{:,2}^1 \partial X_{:,1}^2} & \frac{\partial^2 \underline{R}}{\partial X_{:,2}^1 \partial X_{:,2}^2} & \cdots & \frac{\partial^2 \underline{R}}{\partial X_{:,2}^1 \partial X_{:,M}^2} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \frac{\partial \underline{R}}{\partial X_{:,M}^1} & \frac{\partial^2 \underline{R}}{\partial X_{:,M}^1 \partial X_{:,1}^2} & \frac{\partial^2 \underline{R}}{\partial X_{:,M}^1 \partial X_{:,2}^2} & \cdots & \frac{\partial^2 \underline{R}}{\partial X_{:,M}^1 \partial X_{:,M}^2} \end{bmatrix}$$

$$\frac{\partial \underline{R}}{\partial X_{:,I}^1} = - \left(\frac{\partial \underline{R}}{\partial X_{:,I}^1} \right)^T = - \frac{\partial \underline{R}}{\partial X_{:,I}^2} = \left(\frac{\partial \underline{R}}{\partial X_{:,I}^2} \right)^T$$

$$\frac{\partial^2 \underline{R}}{\partial X_{:,I}^1 \partial X_{:,J}^2} = \left(\frac{\partial^2 \underline{R}}{\partial X_{:,I}^1 \partial X_{:,J}^2} \right)^T = \frac{\partial^2 \underline{R}}{\partial X_{:,J}^1 \partial X_{:,I}^2} = \left(\frac{\partial^2 \underline{R}}{\partial X_{:,J}^1 \partial X_{:,I}^2} \right)^T$$

Here capital I and J are scalar indices for the input dimension (column) of the sample design, $\underline{X}$. Note that for the Gaussian correlation function

$$\frac{\partial^2 R_{j,j}}{\partial X_{j,I}^1 \partial X_{j,I}^2} = 2\theta_I$$

and has units of length⁻². Two of the conditions necessary for a matrix to qualify as a correlation matrix are that all of its elements must be dimensionless and all of its diagonal elements must identically equal one. Since $\underline{R}_{\nabla}$ does not satisfy these requirements, it technically does not qualify as a “correlation matrix.” However, referring to $\underline{R}_{\nabla}$ as such is a small abuse of terminology and allows GEK to use the same naming conventions as Kriging.

A straight-forward implementation of GEK tends to be significantly more accurate than Kriging given the same sample design provided that the

- Derivatives are accurate
- Derivatives are not infinite (or nearly so)

- Function is sufficiently smooth, and
- $\underline{R}_{\nabla}$ is not ill-conditioned (this can be problematic).

If gradients can be obtained cheaply (e.g. by automatic differentiation or adjoint techniques) and the previous conditions are met, GEK also tends to outperform Kriging for the same computational budget. Previous works, such as Dwight[31], state that the direct approach to GEK is significantly better conditioned than the indirect approach. While this is true, (direct) GEK's $\underline{R}_{\nabla}$ matrix can still be, and often is, horribly ill-conditioned compared to Kriging's $\underline{R}$ for the same $\underline{\theta}$ and $\underline{X}$.

In the literature, ill-conditioning is often attributed to the choice of the correlation function. Although a different correlation function may alleviate the ill-conditioning for some problems, the root cause of the ill-conditioning is a poorly spaced sample design. Furthermore, a sufficiently bad sample design could make any interpolatory Kriging model, gradient enhanced or otherwise, ill-conditioned, regardless of the choice of correlation function. This root cause can be addressed directly by discarding points/equations.

Discarding points/equations is conceptually similar to using a Moore-Penrose pseudo inverse of $\underline{R}_{\nabla}$. However, there are important differences. A pseudo inverse handles ill-conditioning by discarding small singular values, which can be interpreted as throwing away the information that is least present while keeping all of what is most frequently duplicated. This causes a Kriging model to not interpolate any of the data points used to construct it while using some information from all rows.

An alternate strategy is to discard additional copies of the information that is most duplicated and keep more of the barely present information. In the context of eigenvalues, this can be described as decreasing the maximum eigenvalue and increasing the minimum eigenvalue by a smaller amount than a pseudo inverse. The result is that the GEK model will exactly fit all of the retained information. This can be achieved using a pivoted Cholesky factorization, such as the one developed by Lucas [66] to determine a reordering $\tilde{\underline{R}}_{\nabla}$ and dropping equations off its end until it tightly meets the constraint on rcond . However, a straight-forward implementation of this is neither efficient nor robust.

In benchmarking tests, Lucas' level 3 pivoted Cholesky implementation was not competitive with the level 3 LAPACK non-pivoted Cholesky in terms of computational efficiency. In some cases, it was an order of magnitude slower. Note that Lucas' level 3 implementation can default to his level 2 implementation and this may explain some of the loss in performance.

More importantly, applying pivoted Cholesky to $\underline{R}_{\nabla}$ tends to sort derivative equations to the top/front and function value equations to the end. This occurred even when $\underline{R}_{\nabla}$ was equilibrated to have ones for all of its diagonal elements. The result was that for many points at least some of the derivative equations were retained while the function values at the same points were discarded. This can (and often did) significantly degrade the accuracy of the GEK predictions. The implication is that derivative equations contain more information than, but are not as reliable as, function value equations.

To address computational efficiency and robustness, Dakota's pivoted Cholesky approach for GEK was modified to:

- Equilibrate $\underline{R}_{\nabla}$ to improve the accuracy of the Cholesky factorization; this is beneficial because $\underline{R}_{\nabla}$ can be poorly scaled. Theorem 4.1 of van der Sluis [94] states that if $\underline{a}$ is a real, symmetric, positive definite n by

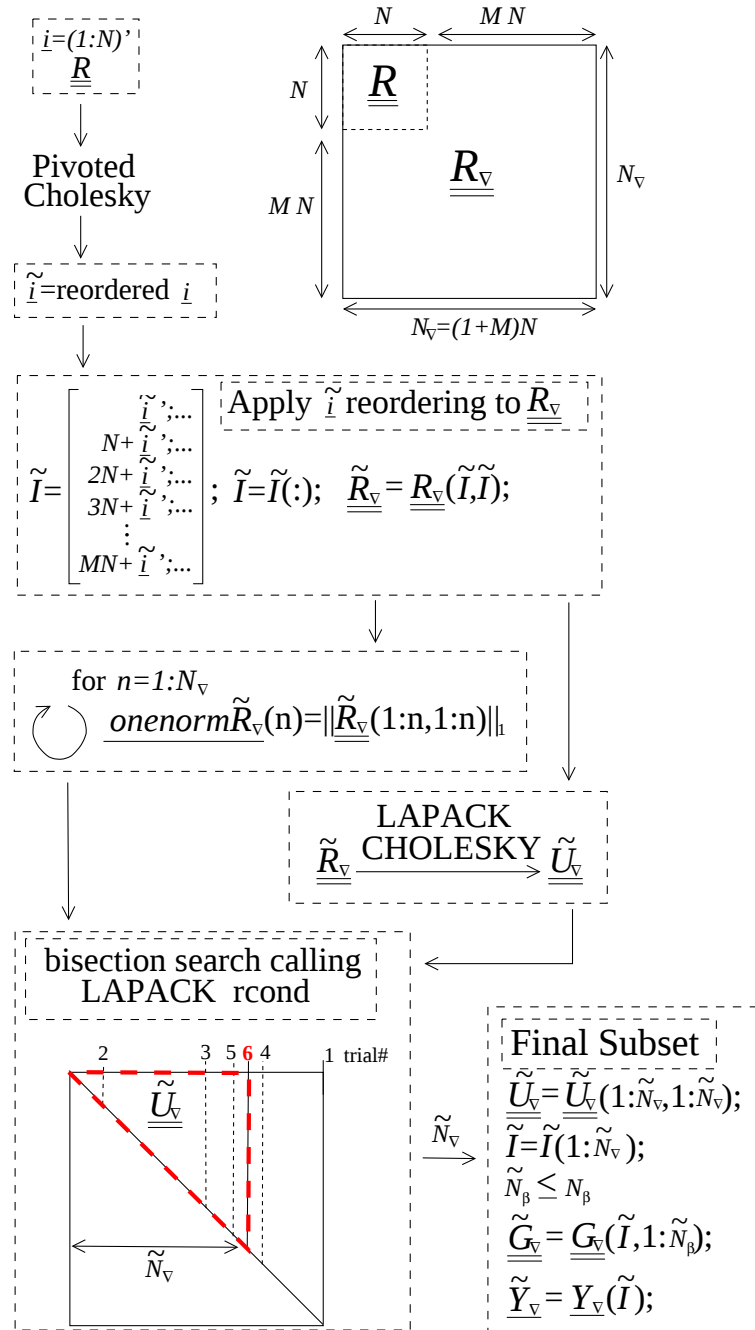


Figure 4.1: A diagram with pseudo code for the pivoted Cholesky algorithm used to select the subset of equations to retain when $\underline{R}_v$ is ill-conditioned. Although it is part of the algorithm, the equilibration of $\underline{R}_v$ is not shown in this figure. The pseudo code uses MATLAB notation.

n matrix and the diagonal matrix $\underline{\underline{\alpha}}$ contains the square roots of the diagonals of $\underline{\underline{a}}$, then the equilibration

$$\underline{\underline{\check{a}}} = \underline{\underline{\alpha}}^{-1} \underline{\underline{a}} \underline{\underline{\alpha}}^{-1},$$

minimizes the 2-norm condition number of $\underline{\underline{\check{a}}}$ (with respect to solving linear systems) over all such symmetric scalings, to within a factor of n . The equilibrated matrix $\underline{\underline{\check{a}}}$ will have all ones on the diagonal.

- Perform pivoted Cholesky on $\underline{\underline{R}}$, instead of $\underline{\underline{R}}_{\nabla}$, to rank points according to how much new information they contain. This ranking was reflected by the ordering of points in $\underline{\underline{\tilde{R}}}$.
- Apply the ordering of points in $\underline{\underline{\tilde{R}}}$ to whole points in $\underline{\underline{R}}_{\nabla}$ to produce $\underline{\underline{\tilde{R}}}_{\nabla}$. Here a whole point means the function value at a point immediately followed by the derivatives at the same point.
- Perform a LAPACK non-pivoted Cholesky on the equilibrated $\underline{\underline{\tilde{R}}}_{\nabla}$ and drop equations off the end until it satisfies the constraint on rcond. LAPACK's rcond estimate requires the 1-norm of the original (reordered) matrix as input so the 1-norms for all possible sizes of $\underline{\underline{\tilde{R}}}_{\nabla}$ are precomputed (using a rank one update approach) and stored prior to the Cholesky factorization. A bisection search is used to efficiently determine the number of equations that need to be retained/discarded. This requires $\text{ceil}(\log 2(N_{\nabla}))$ or fewer evaluations of rcond. These rcond calls are all based off the same Cholesky factorization of $\underline{\underline{\tilde{R}}}_{\nabla}$ but use different numbers of rows/columns, $\tilde{N}_{\nabla}$.

This algorithm is visually depicted in Figure 4.1. Because inverting/factorizing a matrix with n rows and columns requires $\mathcal{O}(n^3)$ flops, the cost to perform pivoted Cholesky on $\underline{\underline{R}}$ will be much less than, i.e. $\mathcal{O}((1+M)^{-3})$, that of $\underline{\underline{R}}_{\nabla}$ when the number of dimensions M is large. It will also likely be negligible compared to the cost of performing LAPACK's non-pivoted Cholesky on $\underline{\underline{\tilde{R}}}_{\nabla}$.

Chapter 5

Surrogate-Based Local Minimization

A generally-constrained nonlinear programming problem takes the form

$$\begin{aligned}
 & \text{minimize} && f(\mathbf{x}) \\
 & \text{subject to} && \mathbf{g}_l \leq \mathbf{g}(\mathbf{x}) \leq \mathbf{g}_u \\
 & && \mathbf{h}(\mathbf{x}) = \mathbf{h}_t \\
 & && \mathbf{x}_l \leq \mathbf{x} \leq \mathbf{x}_u
 \end{aligned} \tag{5.1}$$

where $\mathbf{x} \in \mathbb{R}^n$ is the vector of design variables, and f , $\mathbf{g}$, and $\mathbf{h}$ are the objective function, nonlinear inequality constraints, and nonlinear equality constraints, respectively¹. Individual nonlinear inequality and equality constraints are enumerated using i and j , respectively (e.g., g_i and h_j). The corresponding surrogate-based optimization (SBO) algorithm may be formulated in several ways and applied to either optimization or least-squares calibration problems. In all cases, SBO solves a sequence of k approximate optimization subproblems subject to a trust region constraint Δ^k ; however, many different forms of the surrogate objectives and constraints in the approximate subproblem can be explored. In particular, the subproblem objective may be a surrogate of the original objective or a surrogate of a merit function (most commonly, the Lagrangian or augmented Lagrangian), and the subproblem constraints may be surrogates of the original constraints, linearized approximations of the surrogate constraints, or may be omitted entirely. Each of these combinations is shown in Table 5.1, where black indicates an inappropriate combination, gray indicates an acceptable combination, and blue indicates a common combination.

Initial approaches to nonlinearly-constrained SBO optimized an approximate merit function which incorporated the nonlinear constraints [84, 5]:

$$\begin{aligned}
 & \text{minimize} && \hat{\Phi}^k(\mathbf{x}) \\
 & \text{subject to} && \|\mathbf{x} - \mathbf{x}_c^k\|_\infty \leq \Delta^k
 \end{aligned} \tag{5.2}$$

¹Any linear constraints are not approximated and may be added without modification to all formulations

Table 5.1: SBO approximate subproblem formulations.

	Original Objective	Lagrangian	Augmented Lagrangian
No constraints			TRAL
Linearized constraints		SQP-like	
Original constraints	Direct surrogate		IPTRSAO

where the surrogate merit function is denoted as $\hat{\Phi}(\mathbf{x})$, $\mathbf{x}_c$ is the center point of the trust region, and the trust region is truncated at the global variable bounds as needed. The merit function to approximate was typically chosen to be a standard implementation [95, 74, 49] of the augmented Lagrangian merit function (see Eqs. 5.11–5.12), where the surrogate augmented Lagrangian is constructed from individual surrogate models of the objective and constraints (approximate and assemble, rather than assemble and approximate). In Table 5.1, this corresponds to row 1, column 3, and is known as the trust-region augmented Lagrangian (TRAL) approach. While this approach was provably convergent, convergence rates to constrained minima have been observed to be slowed by the required updating of Lagrange multipliers and penalty parameters [78]. Prior to converging these parameters, SBO iterates did not strictly respect constraint boundaries and were often infeasible. A subsequent approach (IPTRSAO [78]) that sought to directly address this shortcoming added explicit surrogate constraints (row 3, column 3 in Table 5.1):

$$\begin{aligned} & \text{minimize} && \hat{\Phi}^k(\mathbf{x}) \\ & \text{subject to} && \mathbf{g}_l \leq \hat{\mathbf{g}}^k(\mathbf{x}) \leq \mathbf{g}_u \\ & && \hat{\mathbf{h}}^k(\mathbf{x}) = \mathbf{h}_t \\ & && \|\mathbf{x} - \mathbf{x}_c^k\|_\infty \leq \Delta^k. \end{aligned} \quad (5.3)$$

While this approach does address infeasible iterates, it still shares the feature that the surrogate merit function may reflect inaccurate relative weightings of the objective and constraints prior to convergence of the Lagrange multipliers and penalty parameters. That is, one may benefit from more feasible intermediate iterates, but the process may still be slow to converge to optimality. The concept of this approach is similar to that of SQP-like SBO approaches [5] which use linearized constraints:

$$\begin{aligned} & \text{minimize} && \hat{\Phi}^k(\mathbf{x}) \\ & \text{subject to} && \mathbf{g}_l \leq \hat{\mathbf{g}}^k(\mathbf{x}_c^k) + \nabla \hat{\mathbf{g}}^k(\mathbf{x}_c^k)^T (\mathbf{x} - \mathbf{x}_c^k) \leq \mathbf{g}_u \\ & && \hat{\mathbf{h}}^k(\mathbf{x}_c^k) + \nabla \hat{\mathbf{h}}^k(\mathbf{x}_c^k)^T (\mathbf{x} - \mathbf{x}_c^k) = \mathbf{h}_t \\ & && \|\mathbf{x} - \mathbf{x}_c^k\|_\infty \leq \Delta^k. \end{aligned} \quad (5.4)$$

in that the primary concern is minimizing a composite merit function of the objective and constraints, but under the restriction that the original problem constraints may not be wildly violated prior to convergence of Lagrange multiplier estimates. Here, the merit function selection of the Lagrangian function (row 2, column 2 in Table 5.1; see also Eq. 5.10) is most closely related to SQP, which includes the use of first-order Lagrange multiplier updates (Eq. 5.16) that should converge more rapidly near a constrained minimizer than the zeroth-order updates (Eqs. 5.13–5.14) used for the augmented Lagrangian.

All of these previous constrained SBO approaches involve a recasting of the approximate subproblem objective and constraints as a function of the original objective and constraint surrogates. A more direct approach is to use a formulation of:

$$\begin{aligned} & \text{minimize} && \hat{f}^k(\mathbf{x}) \\ & \text{subject to} && \mathbf{g}_l \leq \hat{\mathbf{g}}^k(\mathbf{x}) \leq \mathbf{g}_u \\ & && \hat{\mathbf{h}}^k(\mathbf{x}) = \mathbf{h}_t \\ & && \|\mathbf{x} - \mathbf{x}_c^k\|_\infty \leq \Delta^k \end{aligned} \quad (5.5)$$

This approach has been termed the direct surrogate approach since it optimizes surrogates of the original objective and constraints (row 3, column 1 in Table 5.1) without any recasting. It is attractive both from its simplicity and potential for improved performance, and is the default approach taken in Dakota. Other Dakota defaults include the use of a filter method for iterate acceptance (see Section 5.1), an augmented Lagrangian merit function (see Section 5.2), Lagrangian hard convergence assessment (see Section 5.3), and no constraint relaxation (see Section 5.4).

Table 5.2: Sample trust region ratio logic.

Ratio Value	Surrogate Accuracy	Iterate Acceptance	Trust Region Sizing
$\rho^k \leq 0$	poor	reject step	shrink
$0 < \rho^k \leq 0.25$	marginal	accept step	shrink
$0.25 < \rho^k < 0.75$ or $\rho^k > 1.25$	moderate	accept step	retain
$0.75 \leq \rho^k \leq 1.25$	good	accept step	expand ²

While the formulation of Eq. 5.2 (and others from row 1 in Table 5.1) can suffer from infeasible intermediate iterates and slow convergence to constrained minima, each of the approximate subproblem formulations with explicit constraints (Eqs. 5.3-5.5, and others from rows 2-3 in Table 5.1) can suffer from the lack of a feasible solution within the current trust region. Techniques for dealing with this latter challenge involve some form of constraint relaxation. Homotopy approaches [78, 77] or composite step approaches such as Byrd-Omojokun [75], Celis-Dennis-Tapia [17], or MAESTRO [5] may be used for this purpose (see Section 5.4).

After each of the k iterations in the SBO method, the predicted step is validated by computing $f(\mathbf{x}_*^k)$, $\mathbf{g}(\mathbf{x}_*^k)$, and $\mathbf{h}(\mathbf{x}_*^k)$. One approach forms the trust region ratio ρ^k which measures the ratio of the actual improvement to the improvement predicted by optimization on the surrogate model. When optimizing on an approximate merit function (Eqs. 5.2–5.4), the following ratio is natural to compute

$$\rho^k = \frac{\Phi(\mathbf{x}_c^k) - \Phi(\mathbf{x}_*^k)}{\hat{\Phi}(\mathbf{x}_c^k) - \hat{\Phi}(\mathbf{x}_*^k)}. \quad (5.6)$$

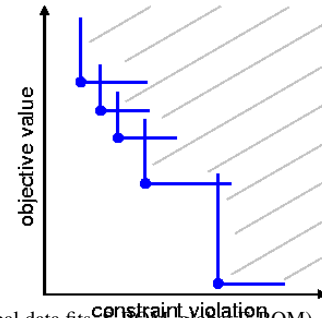
The formulation in Eq. 5.5 may also form a merit function for computing the trust region ratio; however, the omission of this merit function from explicit use in the approximate optimization cycles can lead to synchronization problems with the optimizer.

Once computed, the value for ρ^k can be used to define the step acceptance and the next trust region size Δ^{k+1} using logic similar to that shown in Table 5.2. Typical factors for shrinking and expanding are 0.5 and 2.0, respectively, but these as well as the threshold ratio values are tunable parameters in the algorithm (see Surrogate-Based Method controls in the Dakota Reference Manual [2]). In addition, the use of discrete thresholds is not required, and continuous relationships using adaptive logic can also be explored [102, 103]. Iterate acceptance or rejection completes an SBO cycle, and the cycles are continued until either soft or hard convergence criteria (see Section 5.3) are satisfied.

5.1 Iterate acceptance logic

When a surrogate optimization is completed and the approximate solution has been validated, then the decision must be made to either accept or reject the step. The traditional approach is to base this decision on the value of the trust region ratio, as outlined previously in Table 5.2. An alternate approach is to utilize a filter method [42], which does not require penalty parameters or Lagrange multiplier estimates. The basic idea in a filter method is to apply the concept of Pareto optimality to the objective function and constraint violations and only accept an iterate if it is not dominated by any previous iterate. Mathematically, a new iterate is not dominated if at least one of the following:

$$\text{either } f < f^{(i)} \text{ or } c < c^{(i)} \quad (5.7)$$



²Exception: retain if $\mathbf{x}_*^k$ in trust region interior for design of experiments-based surrogates (global data fits, S-ROM, global E-ROM)

is true for all i in the filter, where c is a selected norm of the constraint violation. This basic description can be augmented with mild requirements to prevent point accumulation and assure convergence, known as a slanting filter [42]. Figure 5.1 illustrates the filter concept, where objective values are plotted against constraint violation for accepted iterates (blue circles) to define the dominated region (denoted by the gray lines). A filter method relaxes the common enforcement of monotonicity in constraint violation reduction and, by allowing more flexibility in acceptable step generation, often allows the algorithm to be more efficient.

The use of a filter method is compatible with any of the SBO formulations in Eqs. 5.2–5.5.

5.2 Merit functions

The merit function $\Phi(\mathbf{x})$ used in Eqs. 5.2–5.4, 5.6 may be selected to be a penalty function, an adaptive penalty function, a Lagrangian function, or an augmented Lagrangian function. In each of these cases, the more flexible inequality and equality constraint formulations with two-sided bounds and targets (Eqs. 5.1, 5.3–5.5), have been converted to a standard form of $\mathbf{g}(\mathbf{x}) \leq 0$ and $\mathbf{h}(\mathbf{x}) = 0$ (in Eqs. 5.8, 5.10–5.16). The active set of inequality constraints is denoted as $\mathbf{g}^+$.

The penalty function employed in this paper uses a quadratic penalty with the penalty schedule linked to SBO iteration number

$$\Phi(\mathbf{x}, r_p) = f(\mathbf{x}) + r_p \mathbf{g}^+(\mathbf{x})^T \mathbf{g}^+(\mathbf{x}) + r_p \mathbf{h}(\mathbf{x})^T \mathbf{h}(\mathbf{x}) \quad (5.8)$$

$$r_p = e^{(k+\text{offset})/10} \quad (5.9)$$

The adaptive penalty function is identical in form to Eq. 5.8, but adapts r_p using monotonic increases in the iteration offset value in order to accept any iterate that reduces the constraint violation.

The Lagrangian merit function is

$$\Phi(\mathbf{x}, \lambda_g, \lambda_h) = f(\mathbf{x}) + \lambda_g^T \mathbf{g}^+(\mathbf{x}) + \lambda_h^T \mathbf{h}(\mathbf{x}) \quad (5.10)$$

for which the Lagrange multiplier estimation is discussed in Section 5.3. Away from the optimum, it is possible for the least squares estimates of the Lagrange multipliers for active constraints to be zero, which equates to omitting the contribution of an active constraint from the merit function. This is undesirable for tracking SBO progress, so usage of the Lagrangian merit function is normally restricted to approximate subproblems and hard convergence assessments.

The augmented Lagrangian employed in this paper follows the sign conventions described in [95]

$$\Phi(\mathbf{x}, \lambda_\psi, \lambda_h, r_p) = f(\mathbf{x}) + \lambda_\psi^T \psi(\mathbf{x}) + r_p \psi(\mathbf{x})^T \psi(\mathbf{x}) + \lambda_h^T \mathbf{h}(\mathbf{x}) + r_p \mathbf{h}(\mathbf{x})^T \mathbf{h}(\mathbf{x}) \quad (5.11)$$

$$\psi_i = \max \left\{ g_i, -\frac{\lambda_{\psi_i}}{2r_p} \right\} \quad (5.12)$$

where $\psi(\mathbf{x})$ is derived from the elimination of slack variables for the inequality constraints. In this case, simple zeroth-order Lagrange multiplier updates may be used:

$$\lambda_\psi^{k+1} = \lambda_\psi^k + 2r_p \psi(\mathbf{x}) \quad (5.13)$$

$$\lambda_h^{k+1} = \lambda_h^k + 2r_p \mathbf{h}(\mathbf{x}) \quad (5.14)$$

The updating of multipliers and penalties is carefully orchestrated [20] to drive reduction in constraint violation of the iterates. The penalty updates can be more conservative than in Eq. 5.9, often using an infrequent application of a constant multiplier rather than a fixed exponential progression.

5.3 Convergence assessment

To terminate the SBO process, hard and soft convergence metrics are monitored. It is preferable for SBO studies to satisfy hard convergence metrics, but this is not always practical (e.g., when gradients are unavailable or unreliable). Therefore, simple soft convergence criteria are also employed which monitor for diminishing returns (relative improvement in the merit function less than a tolerance for some number of consecutive iterations).

To assess hard convergence, one calculates the norm of the projected gradient of a merit function whenever the feasibility tolerance is satisfied. The best merit function for this purpose is the Lagrangian merit function from Eq. 5.10. This requires a least squares estimation for the Lagrange multipliers that best minimize the projected gradient:

$$\nabla_x \Phi(\mathbf{x}, \boldsymbol{\lambda}_g, \boldsymbol{\lambda}_h) = \nabla_x f(\mathbf{x}) + \boldsymbol{\lambda}_g^T \nabla_x \mathbf{g}^+(\mathbf{x}) + \boldsymbol{\lambda}_h^T \nabla_x \mathbf{h}(\mathbf{x}) \quad (5.15)$$

where gradient portions directed into active global variable bounds have been removed. This can be posed as a linear least squares problem for the multipliers:

$$\mathbf{A}\boldsymbol{\lambda} = -\nabla_x f \quad (5.16)$$

where $\mathbf{A}$ is the matrix of active constraint gradients, $\boldsymbol{\lambda}_g$ is constrained to be non-negative, and $\boldsymbol{\lambda}_h$ is unrestricted in sign. To estimate the multipliers using non-negative and bound-constrained linear least squares, the NNLS and BVLS routines [65] from NETLIB are used, respectively.

5.4 Constraint relaxation

The goal of constraint relaxation is to achieve efficiency through the balance of feasibility and optimality when the trust region restrictions prevent the location of feasible solutions to constrained approximate subproblems (Eqs. 5.3-5.5, and other formulations from rows 2-3 in Table 5.1). The SBO algorithm starting from infeasible points will commonly generate iterates which seek to satisfy feasibility conditions without regard to objective reduction [77].

One approach for achieving this balance is to use *relaxed constraints* when iterates are infeasible with respect to the surrogate constraints. We follow Perez, Renaud, and Watson [78], and use a *global homotopy* mapping the relaxed constraints and the surrogate constraints. For formulations in Eqs. 5.3 and 5.5 (and others from row 3 in Table 5.1), the relaxed constraints are defined from

$$\tilde{\mathbf{g}}^k(\mathbf{x}, \tau) = \hat{\mathbf{g}}^k(\mathbf{x}) + (1 - \tau)\mathbf{b}_g \quad (5.17)$$

$$\tilde{\mathbf{h}}^k(\mathbf{x}, \tau) = \hat{\mathbf{h}}^k(\mathbf{x}) + (1 - \tau)\mathbf{b}_h \quad (5.18)$$

For Eq. 5.4 (and others from row 2 in Table 5.1), the original surrogate constraints $\hat{\mathbf{g}}^k(\mathbf{x})$ and $\hat{\mathbf{h}}^k(\mathbf{x})$ in Eqs. 5.17-5.18 are replaced with their linearized forms $(\hat{\mathbf{g}}^k(\mathbf{x}_c^k) + \nabla \hat{\mathbf{g}}^k(\mathbf{x}_c^k)^T(\mathbf{x} - \mathbf{x}_c^k))$ and $(\hat{\mathbf{h}}^k(\mathbf{x}_c^k) + \nabla \hat{\mathbf{h}}^k(\mathbf{x}_c^k)^T(\mathbf{x} - \mathbf{x}_c^k))$, respectively). The approximate subproblem is then reposed using the relaxed constraints as

$$\begin{aligned} & \text{minimize} && \hat{f}^k(\mathbf{x}) \text{ or } \hat{\Phi}^k(\mathbf{x}) \\ & \text{subject to} && \mathbf{g}_l \leq \tilde{\mathbf{g}}^k(\mathbf{x}, \tau^k) \leq \mathbf{g}_u \\ & && \tilde{\mathbf{h}}^k(\mathbf{x}, \tau^k) = \mathbf{h}_t \\ & && \|\mathbf{x} - \mathbf{x}_c^k\|_\infty \leq \Delta^k \end{aligned} \quad (5.19)$$

in place of the corresponding subproblems in Eqs. 5.3-5.5. Alternatively, since the relaxation terms are constants for the k^{th} iteration, it may be more convenient for the implementation to constrain $\hat{\mathbf{g}}^k(\mathbf{x})$ and $\hat{\mathbf{h}}^k(\mathbf{x})$ (or their

linearized forms) subject to relaxed bounds and targets $(\tilde{\mathbf{g}}_l^k, \tilde{\mathbf{g}}_u^k, \tilde{\mathbf{h}}_t^k)$. The parameter τ is the homotopy parameter controlling the extent of the relaxation: when $\tau = 0$, the constraints are fully relaxed, and when $\tau = 1$, the surrogate constraints are recovered. The vectors $\mathbf{b}_g, \mathbf{b}_h$ are chosen so that the starting point, $\mathbf{x}^0$, is feasible with respect to the fully relaxed constraints:

$$\mathbf{g}_l \leq \tilde{\mathbf{g}}^0(\mathbf{x}^0, 0) \leq \mathbf{g}_u \quad (5.20)$$

$$\tilde{\mathbf{h}}^0(\mathbf{x}^0, 0) = \mathbf{h}_t \quad (5.21)$$

At the start of the SBO algorithm, $\tau^0 = 0$ if $\mathbf{x}^0$ is infeasible with respect to the unrelaxed surrogate constraints; otherwise $\tau^0 = 1$ (i.e., no constraint relaxation is used). At the start of the k^{th} SBO iteration where $\tau^{k-1} < 1$, τ^k is determined by solving the subproblem

$$\begin{aligned} & \text{maximize} && \tau^k \\ & \text{subject to} && \mathbf{g}_l \leq \tilde{\mathbf{g}}^k(\mathbf{x}, \tau^k) \leq \mathbf{g}_u \\ & && \tilde{\mathbf{h}}^k(\mathbf{x}, \tau^k) = \mathbf{h}_t \\ & && \|\mathbf{x} - \mathbf{x}_c^k\|_\infty \leq \Delta^k \\ & && \tau^k \geq 0 \end{aligned} \quad (5.22)$$

starting at $(\mathbf{x}_*^{k-1}, \tau^{k-1})$, and then adjusted as follows:

$$\tau^k = \min \{1, \tau^{k-1} + \alpha (\tau_{\max}^k - \tau^{k-1})\} \quad (5.23)$$

The adjustment parameter $0 < \alpha < 1$ is chosen so that the feasible region with respect to the relaxed constraints has positive volume within the trust region. Determining the optimal value for α remains an open question and will be explored in future work.

After τ^k is determined using this procedure, the problem in Eq. 5.19 is solved for $\mathbf{x}_*^k$. If the step is accepted, then the value of τ^k is updated using the current iterate $\mathbf{x}_*^k$ and the validated constraints $\mathbf{g}(\mathbf{x}_*^k)$ and $\mathbf{h}(\mathbf{x}_*^k)$:

$$\tau^k = \min \{1, \min_i \tau_i, \min_j \tau_j\} \quad (5.24)$$

$$\text{where } \tau_i = 1 + \frac{\min\{g_i(\mathbf{x}_*^k) - g_{l_i}, g_{u_i} - g_i(\mathbf{x}_*^k)\}}{b_{g_i}} \quad (5.25)$$

$$\tau_j = 1 - \frac{|h_j(\mathbf{x}_*^k) - h_{t_j}|}{b_{h_j}} \quad (5.26)$$

Figure 5.2 illustrates the SBO algorithm on a two-dimensional problem with one inequality constraint starting from an infeasible point, $\mathbf{x}^0$. The minimizer of the problem is denoted as $\mathbf{x}^*$. Iterates generated using the surrogate constraints are shown in red, where feasibility is achieved first, and then progress is made toward the optimal point. The iterates generated using the relaxed constraints are shown in blue, where a balance of satisfying feasibility and optimality has been achieved, leading to fewer overall SBO iterations.

The behavior illustrated in Fig. 5.2 is an example where using the relaxed constraints over the surrogate constraints may improve the overall performance of the SBO algorithm by reducing the number of iterations performed. This improvement comes at the cost of solving the minimization subproblem in Eq. 5.22, which can be significant in some cases (i.e., when the cost of evaluating $\hat{\mathbf{g}}^k(\mathbf{x})$ and $\hat{\mathbf{h}}^k(\mathbf{x})$ is not negligible, such as with multifidelity or ROM surrogates). As shown in the numerical experiments involving the Barnes

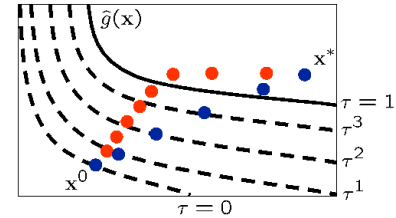


Figure 5.2: Illustration of SBO iterates using surrogate (red) and relaxed (blue) constraints.

problem presented in [78], the directions toward constraint violation reduction and objective function reduction may be in opposing directions. In such cases, the use of the relaxed constraints may result in an *increase* in the overall number of SBO iterations since feasibility must ultimately take precedence.

Chapter 6

Optimization Under Uncertainty (OUU)

6.1 Reliability-Based Design Optimization (RBDO)

Reliability-based design optimization (RBDO) methods are used to perform design optimization accounting for reliability metrics. The reliability analysis capabilities described in Section 1.1 provide a substantial foundation for exploring a variety of gradient-based RBDO formulations. [33] investigated bi-level, fully-analytic bi-level, and first-order sequential RBDO approaches employing underlying first-order reliability assessments. [34] investigated fully-analytic bi-level and second-order sequential RBDO approaches employing underlying second-order reliability assessments. These methods are overviewed in the following sections.

6.1.1 Bi-level RBDO

The simplest and most direct RBDO approach is the bi-level approach in which a full reliability analysis is performed for every optimization function evaluation. This involves a nesting of two distinct levels of optimization within each other, one at the design level and one at the MPP search level.

Since an RBDO problem will typically specify both the $\bar{z}$ level and the $\bar{p}/\bar{\beta}$ level, one can use either the RIA or the PMA formulation for the UQ portion and then constrain the result in the design optimization portion. In particular, RIA reliability analysis maps $\bar{z}$ to p/β , so RIA RBDO constrains p/β :

$$\begin{aligned} & \text{minimize} && f \\ & \text{subject to} && \beta \geq \bar{\beta} \\ & && \text{or } p \leq \bar{p} \end{aligned} \tag{6.1}$$

And PMA reliability analysis maps $\bar{p}/\bar{\beta}$ to z , so PMA RBDO constrains z :

$$\begin{aligned} & \text{minimize} && f \\ & \text{subject to} && z \geq \bar{z} \end{aligned} \tag{6.2}$$

where $z \geq \bar{z}$ is used as the RBDO constraint for a cumulative failure probability (failure defined as $z \leq \bar{z}$) but $z \leq \bar{z}$ would be used as the RBDO constraint for a complementary cumulative failure probability (failure defined as $z \geq \bar{z}$). It is worth noting that Dakota is not limited to these types of inequality-constrained RBDO formulations;

rather, they are convenient examples. Dakota supports general optimization under uncertainty mappings [38] which allow flexible use of statistics within multiple objectives, inequality constraints, and equality constraints.

An important performance enhancement for bi-level methods is the use of sensitivity analysis to analytically compute the design gradients of probability, reliability, and response levels. When design variables are separate from the uncertain variables (i.e., they are not distribution parameters), then the following first-order expressions may be used [56, 62, 6]:

$$\nabla_{\mathbf{d}} z = \nabla_{\mathbf{d}} g \quad (6.3)$$

$$\nabla_{\mathbf{d}} \beta_{cdf} = \frac{1}{\|\nabla_{\mathbf{u}} G\|} \nabla_{\mathbf{d}} g \quad (6.4)$$

$$\nabla_{\mathbf{d}} p_{cdf} = -\phi(-\beta_{cdf}) \nabla_{\mathbf{d}} \beta_{cdf} \quad (6.5)$$

where it is evident from Eqs. 1.10-1.11 that $\nabla_{\mathbf{d}} \beta_{ccdf} = -\nabla_{\mathbf{d}} \beta_{cdf}$ and $\nabla_{\mathbf{d}} p_{ccdf} = -\nabla_{\mathbf{d}} p_{cdf}$. In the case of second-order integrations, Eq. 6.5 must be expanded to include the curvature correction. For Breitung's correction (Eq. 1.38),

$$\nabla_{\mathbf{d}} p_{cdf} = \left[\Phi(-\beta_p) \sum_{i=1}^{n-1} \left(\frac{-\kappa_i}{2(1 + \beta_p \kappa_i)^{\frac{3}{2}}} \prod_{\substack{j=1 \\ j \neq i}}^{n-1} \frac{1}{\sqrt{1 + \beta_p \kappa_j}} \right) - \phi(-\beta_p) \prod_{i=1}^{n-1} \frac{1}{\sqrt{1 + \beta_p \kappa_i}} \right] \nabla_{\mathbf{d}} \beta_{cdf} \quad (6.6)$$

where $\nabla_{\mathbf{d}} \kappa_i$ has been neglected and $\beta_p \geq 0$ (see Section 1.1.2.2). Other approaches assume the curvature correction is nearly independent of the design variables [79], which is equivalent to neglecting the first term in Eq. 6.6.

To capture second-order probability estimates within an RIA RBDO formulation using well-behaved β constraints, a generalized reliability index can be introduced where, similar to Eq. 1.8,

$$\beta_{cdf}^* = -\Phi^{-1}(p_{cdf}) \quad (6.7)$$

for second-order p_{cdf} . This reliability index is no longer equivalent to the magnitude of $\mathbf{u}$, but rather is a convenience metric for capturing the effect of more accurate probability estimates. The corresponding generalized reliability index sensitivity, similar to Eq. 6.5, is

$$\nabla_{\mathbf{d}} \beta_{cdf}^* = -\frac{1}{\phi(-\beta_{cdf}^*)} \nabla_{\mathbf{d}} p_{cdf} \quad (6.8)$$

where $\nabla_{\mathbf{d}} p_{cdf}$ is defined from Eq. 6.6. Even when $\nabla_{\mathbf{d}} g$ is estimated numerically, Eqs. 6.3-6.8 can be used to avoid numerical differencing across full reliability analyses.

When the design variables are distribution parameters of the uncertain variables, $\nabla_{\mathbf{d}} g$ is expanded with the chain rule and Eqs. 6.3 and 6.4 become

$$\nabla_{\mathbf{d}} z = \nabla_{\mathbf{d}} \mathbf{x} \nabla_{\mathbf{x}} g \quad (6.9)$$

$$\nabla_{\mathbf{d}} \beta_{cdf} = \frac{1}{\|\nabla_{\mathbf{u}} G\|} \nabla_{\mathbf{d}} \mathbf{x} \nabla_{\mathbf{x}} g \quad (6.10)$$

where the design Jacobian of the transformation ($\nabla_{\mathbf{d}} \mathbf{x}$) may be obtained analytically for uncorrelated $\mathbf{x}$ or semi-analytically for correlated $\mathbf{x}$ ($\nabla_{\mathbf{d}} \mathbf{L}$ is evaluated numerically) by differentiating Eqs. 1.13 and 1.14 with respect to the distribution parameters. Eqs. 6.5-6.8 remain the same as before. For this design variable case, all required information for the sensitivities is available from the MPP search.

Since Eqs. 6.3-6.10 are derived using the Karush-Kuhn-Tucker optimality conditions for a converged MPP, they are appropriate for RBDO using AMV+, AMV²+, TANA, FORM, and SORM, but not for RBDO using MV-FOSM, MVSOSM, AMV, or AMV².

6.1.2 Sequential/Surrogate-based RBDO

An alternative RBDO approach is the sequential approach, in which additional efficiency is sought through breaking the nested relationship of the MPP and design searches. The general concept is to iterate between optimization and uncertainty quantification, updating the optimization goals based on the most recent probabilistic assessment results. This update may be based on safety factors [100] or other approximations [30].

A particularly effective approach for updating the optimization goals is to use the $p/\beta/z$ sensitivity analysis of Eqs. 6.3-6.10 in combination with local surrogate models [108]. In [33] and [34], first-order and second-order Taylor series approximations were employed within a trust-region model management framework [50] in order to adaptively manage the extent of the approximations and ensure convergence of the RBDO process. Surrogate models were used for both the objective function and the constraints, although the use of constraint surrogates alone is sufficient to remove the nesting.

In particular, RIA trust-region surrogate-based RBDO employs surrogate models of f and p/β within a trust region Δ^k centered at $\mathbf{d}_c$. For first-order surrogates:

$$\begin{aligned} & \text{minimize} && f(\mathbf{d}_c) + \nabla_{\mathbf{d}} f(\mathbf{d}_c)^T (\mathbf{d} - \mathbf{d}_c) \\ & \text{subject to} && \beta(\mathbf{d}_c) + \nabla_{\mathbf{d}} \beta(\mathbf{d}_c)^T (\mathbf{d} - \mathbf{d}_c) \geq \bar{\beta} \\ & && \text{or } p(\mathbf{d}_c) + \nabla_{\mathbf{d}} p(\mathbf{d}_c)^T (\mathbf{d} - \mathbf{d}_c) \leq \bar{p} \\ & && \|\mathbf{d} - \mathbf{d}_c\|_{\infty} \leq \Delta^k \end{aligned} \quad (6.11)$$

and for second-order surrogates:

$$\begin{aligned} & \text{minimize} && f(\mathbf{d}_c) + \nabla_{\mathbf{d}} f(\mathbf{d}_c)^T (\mathbf{d} - \mathbf{d}_c) + \frac{1}{2} (\mathbf{d} - \mathbf{d}_c)^T \nabla_{\mathbf{d}}^2 f(\mathbf{d}_c) (\mathbf{d} - \mathbf{d}_c) \\ & \text{subject to} && \beta(\mathbf{d}_c) + \nabla_{\mathbf{d}} \beta(\mathbf{d}_c)^T (\mathbf{d} - \mathbf{d}_c) + \frac{1}{2} (\mathbf{d} - \mathbf{d}_c)^T \nabla_{\mathbf{d}}^2 \beta(\mathbf{d}_c) (\mathbf{d} - \mathbf{d}_c) \geq \bar{\beta} \\ & && \text{or } p(\mathbf{d}_c) + \nabla_{\mathbf{d}} p(\mathbf{d}_c)^T (\mathbf{d} - \mathbf{d}_c) + \frac{1}{2} (\mathbf{d} - \mathbf{d}_c)^T \nabla_{\mathbf{d}}^2 p(\mathbf{d}_c) (\mathbf{d} - \mathbf{d}_c) \leq \bar{p} \\ & && \|\mathbf{d} - \mathbf{d}_c\|_{\infty} \leq \Delta^k \end{aligned} \quad (6.12)$$

For PMA trust-region surrogate-based RBDO, surrogate models of f and z are employed within a trust region Δ^k centered at $\mathbf{d}_c$. For first-order surrogates:

$$\begin{aligned} & \text{minimize} && f(\mathbf{d}_c) + \nabla_{\mathbf{d}} f(\mathbf{d}_c)^T (\mathbf{d} - \mathbf{d}_c) \\ & \text{subject to} && z(\mathbf{d}_c) + \nabla_{\mathbf{d}} z(\mathbf{d}_c)^T (\mathbf{d} - \mathbf{d}_c) \geq \bar{z} \\ & && \|\mathbf{d} - \mathbf{d}_c\|_{\infty} \leq \Delta^k \end{aligned} \quad (6.13)$$

and for second-order surrogates:

$$\begin{aligned} & \text{minimize} && f(\mathbf{d}_c) + \nabla_{\mathbf{d}} f(\mathbf{d}_c)^T (\mathbf{d} - \mathbf{d}_c) + \frac{1}{2} (\mathbf{d} - \mathbf{d}_c)^T \nabla_{\mathbf{d}}^2 f(\mathbf{d}_c) (\mathbf{d} - \mathbf{d}_c) \\ & \text{subject to} && z(\mathbf{d}_c) + \nabla_{\mathbf{d}} z(\mathbf{d}_c)^T (\mathbf{d} - \mathbf{d}_c) + \frac{1}{2} (\mathbf{d} - \mathbf{d}_c)^T \nabla_{\mathbf{d}}^2 z(\mathbf{d}_c) (\mathbf{d} - \mathbf{d}_c) \geq \bar{z} \\ & && \|\mathbf{d} - \mathbf{d}_c\|_{\infty} \leq \Delta^k \end{aligned} \quad (6.14)$$

where the sense of the z constraint may vary as described previously. The second-order information in Eqs. 6.12 and 6.14 will typically be approximated with quasi-Newton updates.

6.2 Stochastic Expansion-Based Design Optimization (SEBDO)

6.2.1 Stochastic Sensitivity Analysis

Section 2.9 describes sensitivity analysis of the polynomial chaos expansion with respect to the expansion variables. Here we extend this analysis to include sensitivity analysis of probabilistic moments with respect to non-

probabilistic (i.e., design or epistemic uncertain) variables.

6.2.1.1 Local sensitivity analysis: first-order probabilistic expansions

With the introduction of nonprobabilistic variables s (for example, design variables or epistemic uncertain variables), a polynomial chaos expansion only over the probabilistic variables ξ has the functional relationship:

$$R(\xi, s) \cong \sum_{j=0}^P \alpha_j(s) \Psi_j(\xi) \quad (6.15)$$

For computing sensitivities of response mean and variance, the ij indices may be dropped from Eqs. 2.53 and 2.54, simplifying to

$$\mu(s) = \alpha_0(s), \quad \sigma^2(s) = \sum_{k=1}^P \alpha_k^2(s) \langle \Psi_k^2 \rangle \quad (6.16)$$

Sensitivities of Eq. 6.16 with respect to the nonprobabilistic variables are as follows, where independence of s and ξ is assumed:

$$\frac{d\mu}{ds} = \frac{d\alpha_0}{ds} = \left\langle \frac{dR}{ds} \right\rangle \quad (6.17)$$

$$\frac{d\sigma^2}{ds} = \sum_{k=1}^P \langle \Psi_k^2 \rangle \frac{d\alpha_k^2}{ds} = 2 \sum_{k=1}^P \alpha_k \left\langle \frac{dR}{ds}, \Psi_k \right\rangle \quad (6.18)$$

where

$$\frac{d\alpha_k}{ds} = \frac{\langle \frac{dR}{ds}, \Psi_k \rangle}{\langle \Psi_k^2 \rangle} \quad (6.19)$$

has been used. Due to independence, the coefficients calculated in Eq. 6.19 may be interpreted as either the derivatives of the expectations or the expectations of the derivatives, or more precisely, the nonprobabilistic sensitivities of the chaos coefficients for the response expansion or the chaos coefficients of an expansion for the nonprobabilistic sensitivities of the response. The evaluation of integrals involving $\frac{dR}{ds}$ extends the data requirements for the PCE approach to include response sensitivities at each of the sampled points. The resulting expansions are valid only for a particular set of nonprobabilistic variables and must be recalculated each time the nonprobabilistic variables are modified.

Similarly for stochastic collocation,

$$R(\xi, s) \cong \sum_{k=1}^{N_p} r_k(s) L_k(\xi) \quad (6.20)$$

leads to

$$\mu(s) = \sum_{k=1}^{N_p} r_k(s) w_k, \quad \sigma^2(s) = \sum_{k=1}^{N_p} r_k^2(s) w_k - \mu^2(s) \quad (6.21)$$

$$\frac{d\mu}{ds} = \sum_{k=1}^{N_p} w_k \frac{dr_k}{ds} \quad (6.22)$$

$$\frac{d\sigma^2}{ds} = \sum_{k=1}^{N_p} 2w_k r_k \frac{dr_k}{ds} - 2\mu \frac{d\mu}{ds} = \sum_{k=1}^{N_p} 2w_k (r_k - \mu) \frac{dr_k}{ds} \quad (6.23)$$

6.2.1.2 Local sensitivity analysis: zeroth-order combined expansions

Alternatively, a stochastic expansion can be formed over both ξ and s . Assuming a bounded design domain $s_L \leq s \leq s_U$ (with no implied probability content), a Legendre chaos basis would be appropriate for each of the dimensions in s within a polynomial chaos expansion.

$$R(\xi, s) \cong \sum_{j=0}^P \alpha_j \Psi_j(\xi, s) \quad (6.24)$$

In this case, design sensitivities for the mean and variance do not require response sensitivity data, but this comes at the cost of forming the PCE over additional dimensions. For this combined variable expansion, the mean and variance are evaluated by performing the expectations over only the probabilistic expansion variables, which eliminates the polynomial dependence on ξ , leaving behind the desired polynomial dependence of the moments on s :

$$\mu_R(s) = \sum_{j=0}^P \alpha_j \langle \Psi_j(\xi, s) \rangle_{\xi} \quad (6.25)$$

$$\sigma_R^2(s) = \sum_{j=0}^P \sum_{k=0}^P \alpha_j \alpha_k \langle \Psi_j(\xi, s) \Psi_k(\xi, s) \rangle_{\xi} - \mu_R^2(s) \quad (6.26)$$

The remaining polynomials may then be differentiated with respect to s . In this approach, the combined PCE is valid for the full design variable range ($s_L \leq s \leq s_U$) and does not need to be updated for each change in nonprobabilistic variables, although adaptive localization techniques (i.e., trust region model management approaches) can be employed when improved local accuracy of the sensitivities is required.

Similarly for stochastic collocation,

$$R(\xi, s) \cong \sum_{j=1}^{N_p} r_j L_j(\xi, s) \quad (6.27)$$

leads to

$$\mu_R(s) = \sum_{j=1}^{N_p} r_j \langle L_j(\xi, s) \rangle_{\xi} \quad (6.28)$$

$$\sigma_R^2(s) = \sum_{j=1}^{N_p} \sum_{k=1}^{N_p} r_j r_k \langle L_j(\xi, s) L_k(\xi, s) \rangle_{\xi} - \mu_R^2(s) \quad (6.29)$$

where the remaining polynomials not eliminated by the expectation over ξ are again differentiated with respect to s .

6.2.1.3 Inputs and outputs

There are two types of nonprobabilistic variables for which sensitivities must be calculated: “augmented,” where the nonprobabilistic variables are separate from and augment the probabilistic variables, and “inserted,” where the nonprobabilistic variables define distribution parameters for the probabilistic variables. Any inserted nonprobabilistic variable sensitivities must be handled using Eqs. 6.17-6.18 and Eqs. 6.22-6.23 where $\frac{dR}{ds}$ is calculated as

$\frac{dR}{dx} \frac{dx}{ds}$ and $\frac{dx}{ds}$ is the Jacobian of the variable transformation $\mathbf{x} = T^{-1}(\boldsymbol{\xi})$ with respect to the inserted nonprobabilistic variables. In addition, parameterized polynomials (generalized Gauss-Laguerre, Jacobi, and numerically-generated polynomials) may introduce a $\frac{d\Psi}{ds}$ or $\frac{dL}{ds}$ dependence for inserted s that will introduce additional terms in the sensitivity expressions.

While moment sensitivities directly enable robust design optimization and interval estimation formulations which seek to control or bound response variance, control or bounding of reliability requires sensitivities of tail statistics. In this work, the sensitivity of simple moment-based approximations to cumulative distribution function (CDF) and complementary cumulative distribution function (CCDF) mappings (Eqs. 1.3–1.4) are employed for this purpose, such that it is straightforward to form approximate design sensitivities of reliability index β (forward reliability mapping $\bar{z} \rightarrow \beta$) or response level z (inverse reliability mapping $\beta \rightarrow z$) from the moment design sensitivities and the specified levels $\bar{\beta}$ or $\bar{z}$.

6.2.2 Optimization Formulations

Given the capability to compute analytic statistics of the response along with design sensitivities of these statistics, Dakota supports bi-level, sequential, and multifidelity approaches for optimization under uncertainty (OUU). The latter two approaches apply surrogate modeling approaches (data fits and multifidelity modeling) to the uncertainty analysis and then apply trust region model management to the optimization process.

6.2.2.1 Bi-level SEBDO

The simplest and most direct approach is to employ these analytic statistics and their design derivatives from Section 6.2.1 directly within an optimization loop. This approach is known as bi-level OUU, since there is an inner level uncertainty analysis nested within an outer level optimization.

Consider the common reliability-based design example of a deterministic objective function with a reliability constraint:

$$\begin{aligned} & \text{minimize} && f \\ & \text{subject to} && \beta \geq \bar{\beta} \end{aligned} \tag{6.30}$$

where β is computed relative to a prescribed threshold response value $\bar{z}$ (e.g., a failure threshold) and is constrained by a prescribed reliability level $\bar{\beta}$ (minimum allowable reliability in the design), and is either a CDF or CCDF index depending on the definition of the failure domain (i.e., defined from whether the associated failure probability is cumulative, $p(g \leq \bar{z})$, or complementary cumulative, $p(g > \bar{z})$).

Another common example is robust design in which the constraint enforcing a reliability lower-bound has been replaced with a constraint enforcing a variance upper-bound $\bar{\sigma}^2$ (maximum allowable variance in the design):

$$\begin{aligned} & \text{minimize} && f \\ & \text{subject to} && \sigma^2 \leq \bar{\sigma}^2 \end{aligned} \tag{6.31}$$

Solving these problems using a bi-level approach involves computing β and $\frac{d\beta}{ds}$ for Eq. 6.30 or σ^2 and $\frac{d\sigma^2}{ds}$ for Eq. 6.31 for each set of design variables s passed from the optimizer. This approach is supported for both probabilistic and combined expansions using PCE and SC.

6.2.2.2 Sequential/Surrogate-Based SEBDO

An alternative OUU approach is the sequential approach, in which additional efficiency is sought through breaking the nested relationship of the UQ and optimization loops. The general concept is to iterate between optimization and uncertainty quantification, updating the optimization goals based on the most recent uncertainty assessment results. This approach is common with the reliability methods community, for which the updating strategy may be based on safety factors [100] or other approximations [30].

A particularly effective approach for updating the optimization goals is to use data fit surrogate models, and in particular, local Taylor series models allow direct insertion of stochastic sensitivity analysis capabilities. In Ref. [33], first-order Taylor series approximations were explored, and in Ref. [34], second-order Taylor series approximations are investigated. In both cases, a trust-region model management framework [36] is used to adaptively manage the extent of the approximations and ensure convergence of the OUU process. Surrogate models are used for both the objective and the constraint functions, although the use of surrogates is only required for the functions containing statistical results; deterministic functions may remain explicit is desired.

In particular, trust-region surrogate-based optimization for reliability-based design employs surrogate models of f and β within a trust region Δ^k centered at $\mathbf{s}_c$:

$$\begin{aligned} & \text{minimize} && f(\mathbf{s}_c) + \nabla_s f(\mathbf{s}_c)^T (\mathbf{s} - \mathbf{s}_c) \\ & \text{subject to} && \beta(\mathbf{s}_c) + \nabla_s \beta(\mathbf{s}_c)^T (\mathbf{s} - \mathbf{s}_c) \geq \bar{\beta} \\ & && \|\mathbf{s} - \mathbf{s}_c\|_\infty \leq \Delta^k \end{aligned} \quad (6.32)$$

and trust-region surrogate-based optimization for robust design employs surrogate models of f and σ^2 within a trust region Δ^k centered at $\mathbf{s}_c$:

$$\begin{aligned} & \text{minimize} && f(\mathbf{s}_c) + \nabla_s f(\mathbf{s}_c)^T (\mathbf{s} - \mathbf{s}_c) \\ & \text{subject to} && \sigma^2(\mathbf{s}_c) + \nabla_s \sigma^2(\mathbf{s}_c)^T (\mathbf{s} - \mathbf{s}_c) \leq \bar{\sigma}^2 \\ & && \|\mathbf{s} - \mathbf{s}_c\|_\infty \leq \Delta^k \end{aligned} \quad (6.33)$$

Second-order local surrogates may also be employed, where the Hessians are typically approximated from an accumulation of curvature information using quasi-Newton updates [74] such as Broyden-Fletcher-Goldfarb-Shanno (BFGS, Eq. 1.42) or symmetric rank one (SR1, Eq. 1.43). The sequential approach is available for probabilistic expansions using PCE and SC.

6.2.2.3 Multifidelity SEBDO

The multifidelity OUU approach is another trust-region surrogate-based approach. Instead of the surrogate UQ model being a simple data fit (in particular, first-/second-order Taylor series model) of the truth UQ model results, distinct UQ models of differing fidelity are now employed. This differing UQ fidelity could stem from the fidelity of the underlying simulation model, the fidelity of the UQ algorithm, or both. In this section, we focus on the fidelity of the UQ algorithm. For reliability methods, this could entail varying fidelity in approximating assumptions (e.g., Mean Value for low fidelity, SORM for high fidelity), and for stochastic expansion methods, it could involve differences in selected levels of p and h refinement.

Here, we define UQ fidelity as point-wise accuracy in the design space and take the high fidelity truth model to be the probabilistic expansion PCE/SC model, with validity only at a single design point. The low fidelity model, whose validity over the design space will be adaptively controlled, will be either the combined expansion PCE/SC model, with validity over a range of design parameters, or the MVFOSM reliability method, with validity only at a single design point. The combined expansion low fidelity approach will span the current trust region

of the design space and will be reconstructed for each new trust region. Trust region adaptation will ensure that the combined expansion approach remains sufficiently accurate for design purposes. By taking advantage of the design space spanning, one can eliminate the cost of multiple low fidelity UQ analyses within the trust region, with fallback to the greater accuracy and higher expense of the probabilistic expansion approach when needed. The MVFOSM low fidelity approximation must be reformed for each change in design variables, but it only requires a single evaluation of a response function and its derivative to approximate the response mean and variance from the input mean and covariance (Eqs. 1.1–1.2) from which forward/inverse CDF/CCDF reliability mappings can be generated using Eqs. 1.3–1.4. This is the least expensive UQ option, but its limited accuracy¹ may dictate the use of small trust regions, resulting in greater iterations to convergence. The expense of optimizing a combined expansion, on the other hand, is not significantly less than that of optimizing the high fidelity UQ model, but its representation of global trends should allow the use of larger trust regions, resulting in reduced iterations to convergence. The design derivatives of each of the PCE/SC expansion models provide the necessary data to correct the low fidelity model to first-order consistency with the high fidelity model at the center of each trust region, ensuring convergence of the multifidelity optimization process to the high fidelity optimum. Design derivatives of the MVFOSM statistics are currently evaluated numerically using forward finite differences.

Multifidelity optimization for reliability-based design can be formulated as:

$$\begin{aligned} & \text{minimize} && f(\mathbf{s}) \\ & \text{subject to} && \hat{\beta}_{hi}(\mathbf{s}) \geq \bar{\beta} \\ & && \|\mathbf{s} - \mathbf{s}_c\|_{\infty} \leq \Delta^k \end{aligned} \quad (6.34)$$

and multifidelity optimization for robust design can be formulated as:

$$\begin{aligned} & \text{minimize} && f(\mathbf{s}) \\ & \text{subject to} && \hat{\sigma}_{hi}^2(\mathbf{s}) \leq \bar{\sigma}^2 \\ & && \|\mathbf{s} - \mathbf{s}_c\|_{\infty} \leq \Delta^k \end{aligned} \quad (6.35)$$

where the deterministic objective function is not approximated and $\hat{\beta}_{hi}$ and $\hat{\sigma}_{hi}^2$ are the approximated high-fidelity UQ results resulting from correction of the low-fidelity UQ results. In the case of an additive correction function:

$$\hat{\beta}_{hi}(\mathbf{s}) = \beta_{lo}(\mathbf{s}) + \alpha_{\beta}(\mathbf{s}) \quad (6.36)$$

$$\hat{\sigma}_{hi}^2(\mathbf{s}) = \sigma_{lo}^2(\mathbf{s}) + \alpha_{\sigma^2}(\mathbf{s}) \quad (6.37)$$

where correction functions $\alpha(\mathbf{s})$ enforcing first-order consistency [37] are typically employed. Quasi-second-order correction functions [37] can also be employed, but care must be taken due to the different rates of curvature accumulation between the low and high fidelity models. In particular, since the low fidelity model is evaluated more frequently than the high fidelity model, it accumulates curvature information more quickly, such that enforcing quasi-second-order consistency with the high fidelity model can be detrimental in the initial iterations of the algorithm². Instead, this consistency should only be enforced when sufficient high fidelity curvature information has been accumulated (e.g., after n rank one updates).

¹MVFOSM is exact for linear functions with Gaussian inputs, but quickly degrades for nonlinear and/or non-Gaussian.

²Analytic and numerical Hessians, when available, are instantaneous with no accumulation rate concerns.

Bibliography

- [1] M. Abramowitz and I. A. Stegun. *Handbook of Mathematical Functions with Formulas, Graphs, and Mathematical Tables*. Dover, New York, 1965. 22, 30, 38
- [2] B. M. Adams, L. E. Bauman, W. J. Bohnhoff, K. R. Dalbey, J. P. Eddy, M. S. Ebeida, M. S. Eldred, P. D. Hough, K. T. Hu, J. D. Jakeman, L. P. Swiler, Stephens J. A., D. M. Vigil, and T. M. Wildey. Dakota, a multilevel parallel object-oriented framework for design optimization, parameter estimation, uncertainty quantification, and sensitivity analysis: Version 6.0 reference manual. Technical Report SAND2014-xxxx, Sandia National Laboratories, Albuquerque, NM, Updated May 2014. Available online from <http://dakota.sandia.gov/documentation.html>. 34, 40, 59
- [3] B. M. Adams, L. E. Bauman, W. J. Bohnhoff, K. R. Dalbey, J. P. Eddy, M. S. Ebeida, M. S. Eldred, P. D. Hough, K. T. Hu, J. D. Jakeman, L. P. Swiler, Stephens J. A., D. M. Vigil, and T. M. Wildey. Dakota, a multilevel parallel object-oriented framework for design optimization, parameter estimation, uncertainty quantification, and sensitivity analysis: Version 6.0 users manual. Technical Report SAND2014-xxxx, Sandia National Laboratories, Albuquerque, NM, Updated May 2014. Available online from <http://dakota.sandia.gov/documentation.html>. 38, 42
- [4] N. Agarwal and N.R. Aluru. A domain adaptive stochastic collocation approach for analysis of MEMS under uncertainties. *Journal of Computational Physics*, 228:7662–7688, 2009. 24
- [5] N. M. Alexandrov, R. M. Lewis, C. R. Gumbert, L. L. Green, and P. A. Newman. Optimization with variable-fidelity models applied to wing design. In *Proceedings of the 38th Aerospace Sciences Meeting and Exhibit*, Reno, NV, 2000. AIAA Paper 2000-0841. 57, 58, 59
- [6] M. Allen and K. Maute. Reliability-based design optimization of aeroelastic structures. *Struct. Multidiscip. O.*, 27:228–242, 2004. 66
- [7] R. Askey and J. Wilson. Some basic hypergeometric polynomials that generalize jacobi polynomials. In *Mem. Amer. Math. Soc.* 319, Providence, RI, 1985. AMS. 21
- [8] V. Barthelmann, E. Novak, and K. Ritter. High dimensional polynomial interpolation on sparse grids. *Adv. Comput. Math.*, 12(4):273–288, 2000. Multivariate polynomial interpolation. 32
- [9] C. Berg, J. Mateu, and E. Porcu. The dagum family of isotropic correlation functions. *Bernoulli*, 14(4):1134–1149, 2008. 49
- [10] J.-P. Berrut and L. N. Trefethen. Barycentric lagrange interpolation. *SIAM Review*, 46(3):501–517, 2004. 23
- [11] Graud Blatman and Bruno Sudret. Adaptive sparse polynomial chaos expansion based on least angle regression. *Journal of Computational Physics*, 230(6):2345 – 2367, 2011. 34

- [12] G. E. P. Box and D. R. Cox. An analysis of transformations. *J. Royal Stat. Soc., Series B*, 26:211–252, 1964. 10, 30
- [13] S. Boyd and L. Vandenberghe. *Convex Optimization*. Cambridge University Press, March 2004. 34
- [14] K. Breitung. Asymptotic approximation for multinormal integrals. *J. Eng. Mech., ASCE*, 110(3):357–366, 1984. 14
- [15] J. Burkardt. The “combining coefficient” for anisotropic sparse grids. Technical report, Virginia Tech, Blacksburg, VA, 2009. http://people.sc.fsu.edu/~jburkardt/presentations/sgmga_coefficient.pdf. 33
- [16] J. Burkardt. Computing the hermite interpolation polynomial. Technical report, Florida State University, Gainesville, FL, 2009. http://people.sc.fsu.edu/~jburkardt/presentations/hermite_interpolant.pdf. 23
- [17] M. R. Celis, J. .E. Dennis, and R. .A. Tapia. A trust region strategy for nonlinear equality constrained optimization. In P. .T. Boggs, R. H. Byrd, and R. B. Schnabel, editors, *Numerical Optimization 1984*, pages 71–82. SIAM, Philadelphia, USA, 1985. 59
- [18] S. Chen, D. Donoho, and M. Saunders. Atomic decomposition by basis pursuit. *SIAM Rev.*, 43(1):129–159, January 2001. 34
- [19] X. Chen and N.C. Lind. Fast probability integration by three-parameter normal tail approximation. *Struct. Saf.*, 1:269–276, 1983. 10
- [20] A. R. Conn, N. I. M. Gould, and P. L. Toint. *Trust-Region Methods*. MPS-SIAM Series on Optimization, SIAM-MPS, Philadelphia, 2000. 60
- [21] P. G. Constantine, M. S. Eldred, and E. T. Phipps. Sparse pseudospectral approximation method. *Computer Methods in Applied Mechanics and Engineering*, 229–232:1–12, July 2012. 33
- [22] P. G. Constantine, D. F. Gleich, and G. Iaccarino. Spectral methods for parameterized matrix equations. *SIAM Journal on Matrix Analysis and Applications*, 31(5):2681–2699, 2010. 32
- [23] N. Cressie. *Statistics of Spatial Data*. John Wiley and Sons, New York, 1991. 17, 47
- [24] G. Davis, S. Mallat, and M. Avellaneda. Adaptive greedy approximations. *Constructive Approximation*, 13:57–98, 1997. 10.1007/BF02678430. 35
- [25] G.J. Davis and M.D. Morris. Six factors which affect the condition number of matrices associated with kriging. *Mathematical geology*, 29(5):669–683, 1997. 49
- [26] A. Der Kiureghian and P. L. Liu. Structural reliability under incomplete information. *J. Eng. Mech., ASCE*, 112(EM-1):85–104, 1986. 10, 11, 22, 30
- [27] A. Dey and S. Mahadevan. Ductile structural system reliability analysis using adaptive importance sampling. *Structural Safety*, 20:137–154, 1998. 16
- [28] D.L. Donoho and Y. Tsaig. Fast solution of ℓ_1 -norm minimization problems when the solution may be sparse. *Information Theory, IEEE Transactions on*, 54(11):4789–4812, nov. 2008. 36
- [29] A. . Doostan and H. Owhadi. A non-adapted sparse approximation of PDEs with stochastic inputs. *Journal of Computational Physics*, 230(8):3015 – 3034, 2011. 34

- [30] X. Du and W. Chen. Sequential optimization and reliability assessment method for efficient probabilistic design. *J. Mech. Design*, 126:225–233, 2004. [67](#), [71](#)
- [31] R.P. Dwight and Z. Han. Efficient uncertainty quantification using gradient-enhanced kriging. In *11th AIAA Non-Deterministic Approaches Conference*, Reston, VA, USA, May 2009. American Institute of Aeronautics and Astronautics. [54](#)
- [32] B. Efron, T. Hastie, I. Johnstone, and R. Tibshirani. Least angle regression. *Annals of Statistics*, 32:407–499, 2004. [35](#), [36](#)
- [33] M. S. Eldred, H. Agarwal, V. M. Perez, S. F. Wojtkiewicz, Jr., and J. E. Renaud. Investigation of reliability method formulations in DAKOTA/UQ. *Structure & Infrastructure Engineering: Maintenance, Management, Life-Cycle Design & Performance*, 3(3):199–213, 2007. [12](#), [15](#), [65](#), [67](#), [71](#)
- [34] M. S. Eldred and B. J. Bichon. Second-order reliability formulations in DAKOTA/UQ. In *Proceedings of the 47th AIAA/ASME/ASCE/AHS/ASC Structures, Structural Dynamics and Materials Conference*, number AIAA-2006-1828, Newport, RI, May 1–4 2006. [12](#), [15](#), [65](#), [67](#), [71](#)
- [35] M. S. Eldred and J. Burkardt. Comparison of non-intrusive polynomial chaos and stochastic collocation methods for uncertainty quantification. In *Proceedings of the 47th AIAA Aerospace Sciences Meeting and Exhibit*, number AIAA-2009-0976, Orlando, FL, January 5–8, 2009. [31](#), [32](#), [33](#)
- [36] M. S. Eldred and D. M. Dunlavy. Formulations for surrogate-based optimization with data fit, multifidelity, and reduced-order models. In *Proceedings of the 11th AIAA/ISSMO Multidisciplinary Analysis and Optimization Conference*, number AIAA-2006-7117, Portsmouth, VA, September 6–8 2006. [71](#)
- [37] M. S. Eldred, A. A. Giunta, and S. S. Collis. Second-order corrections for surrogate-based optimization with model hierarchies. In *Proceedings of the 10th AIAA/ISSMO Multidisciplinary Analysis and Optimization Conference*, Albany, NY, Aug. 30–Sept. 1, 2004. AIAA Paper 2004-4457. [72](#)
- [38] M. S. Eldred, A. A. Giunta, S. F. Wojtkiewicz, Jr., and T. G. Trucano. Formulations for surrogate-based optimization under uncertainty. In *Proc. 9th AIAA/ISSMO Symposium on Multidisciplinary Analysis and Optimization*, number AIAA-2002-5585, Atlanta, GA, September 4–6, 2002. [66](#)
- [39] M. S. Eldred, C. G. Webster, and P. Constantine. Evaluation of non-intrusive approaches for wiener-asky generalized polynomial chaos. In *Proceedings of the 10th AIAA Non-Deterministic Approaches Conference*, number AIAA-2008-1892, Schaumburg, IL, April 7–10 2008. [22](#)
- [40] M.S. Eldred, B.M. Adams, D.M. Gay, L.P. Swiler, K. Haskell, W.J. Bohnhoff, J.P. Eddy, W.E. Hart, J.P. Watson, J.D. Griffin, P.D. Hough, T.G. Kolda, P.J. Williams, and M.L. Martinez-Canales. Dakota version 4.1 users manual. Sandia Technical Report SAND2006-6337, Sandia National Laboratories, Albuquerque, NM, 2007. [49](#)
- [41] G. M. Fadel, M. F. Riley, and J.-F. M. Barthelemy. Two point exponential approximation method for structural optimization. *Structural Optimization*, 2(2):117–124, 1990. [13](#)
- [42] R. Fletcher, S. Leyffer, and P. L. Toint. On the global convergence of a filter-SQP algorithm. *SIAM J. Optim.*, 13(1):44–59, 2002. [59](#), [60](#)
- [43] P. Frauenfelder, C. Schwab, and R. A. Todor. Finite elements for elliptic problems with stochastic coefficients. *Comput. Methods Appl. Mech. Engrg.*, 194(2-5):205–228, 2005. [32](#)
- [44] J. Gablonsky. Direct version 2.0 userguide technical report. Technical Report CRSC-TR01-08, North Carolina State University, Center for Research in Scientific Computation, Raleigh, NC, 2001. [19](#)

- [45] W. Gautschi. *Orthogonal Polynomials: Computation and Approximation*. Oxford University Press, New York, 2004. [22](#)
- [46] T. Gerstner and M. Griebel. Numerical integration using sparse grids. *Numer. Algorithms*, 18(3-4):209–232, 1998. [32](#)
- [47] T. Gerstner and M. Griebel. Dimension-adaptive tensor-product quadrature. *Computing*, 71(1):65–87, 2003. [41](#)
- [48] P. E. Gill, W. Murray, M. A. Saunders, and M. H. Wright. User’s guide for NPSOL (Version 4.0): A Fortran package for nonlinear programming. Technical Report TR SOL-86-2, System Optimization Laboratory, Stanford University, Stanford, CA, 1986. [15](#)
- [49] P. E. Gill, W. Murray, and M. H. Wright. *Practical Optimization*. Academic Press, San Diego, CA, 1981. [58](#)
- [50] A. A. Giunta and M. S. Eldred. Implementation of a trust region model management strategy in the DAKOTA optimization toolkit. In *Proc. 8th AIAA/USAF/NASA/ISSMO Symposium on Multidisciplinary Analysis and Optimization*, number AIAA-2000-4935, Long Beach, CA, September 6–8, 2000. [67](#)
- [51] T. Gneiting, M.G. Genton, and P. Guttorp. Geostatistical space-time models, stationarity, separability and full symmetry. In B. Finkenstadt, L. Held, and V. Isham, editors, *Statistical Methods for Spatio-Temporal Systems*, chapter 4, pages 151–172. Boca Raton: Chapman & Hall/CRC, 2007. [49](#)
- [52] G. H. Golub and J. H. Welsch. Calculation of gauss quadrature rules. *Mathematics of Computation*, 23(106):221–230, 1969. [22](#)
- [53] A. Haldar and S. Mahadevan. *Probability, Reliability, and Statistical Methods in Engineering Design*. Wiley, New York, 2000. [9](#), [10](#), [15](#)
- [54] T. Hastie, R. Tibshirani, and J.H. Friedman. *The elements of statistical learning: data mining, inference, and prediction: with 200 full-color illustrations*. New York: Springer-Verlag, 2001. [36](#)
- [55] N. J. Higham. The numerical stability of barycentric lagrange interpolation. *IMA Journal of Numerical Analysis*, 24(4):547–556, 2004. [23](#)
- [56] M. Hohenbichler and R. Rackwitz. Sensitivity and importance measures in structural reliability. *Civil Eng. Syst.*, 3:203–209, 1986. [66](#)
- [57] M. Hohenbichler and R. Rackwitz. Improvement of second-order reliability estimates by importance sampling. *J. Eng. Mech., ASCE*, 114(12):2195–2199, 1988. [14](#)
- [58] H.P. Hong. Simple approximations for improving second-order reliability estimates. *J. Eng. Mech., ASCE*, 125(5):592–595, 1999. [14](#)
- [59] S. Hosder, R. W. Walters, and M. Balch. Efficient sampling for non-intrusive polynomial chaos applications with multiple uncertain input variables. In *Proceedings of the 48th AIAA/ASME/ASCE/AHS/ASC Structures, Structural Dynamics, and Materials Conference*, number AIAA-2007-1939, Honolulu, HI, April 23–26, 2007. [34](#)
- [60] D. Huang, T. T. Allen, W. I. Notz, and N. Zeng. Global optimization of stochastic black-box systems via sequential kriging meta-models. *Journal of Global Optimization*, 34:441–466, 2006. [16](#), [19](#)
- [61] D. Jones, M. Schonlau, and W. Welch. Efficient global optimization of expensive black-box functions. *Journal of Global Optimization*, 13:455–492, 1998. [16](#), [18](#), [19](#)

- [62] A. Karamchandani and C. A. Cornell. Sensitivity estimation within first and second order reliability methods. *Struct. Saf.*, 11:95–107, 1992. [66](#)
- [63] A. Klimke and B. Wohlmuth. Algorithm 847: spinterp: Piecewise multilinear hierarchical sparse grid interpolation in matlab. *ACM Transactions on Mathematical Software*, 31(4):561–579, 2005. [24](#)
- [64] W. A. Klimke. *Uncertainty Modeling using Fuzzy Arithmetic and Sparse Grids*. PhD thesis, Universität Stuttgart, Stuttgart, Germany, 2005. [28](#)
- [65] C. L. Lawson and R. J. Hanson. *Solving Least Squares Problems*. Prentice–Hall, 1974. [61](#)
- [66] C. Lucas. Lapack-style codes for level 2 and 3 pivoted cholesky factorizations. Numerical Analysis Reports 442, Manchester Centre for Computational Mathematics, Manchester, England, February 2004. LAPACK Working Note 161. [54](#)
- [67] J.D. Martin. Computational improvements to estimating kriging metamodel parameters. *Journal of Mechanical Design*, 131:084501, 2009. [49](#)
- [68] G. Matheron. *The theory of regionalized variables and its applications*. Les Cahiers du Centre de morphologie mathématique de Fontainebleau. École nationale supérieure des mines, 1971. [48](#)
- [69] J. C. Meza, R. A. Oliva, P. D. Hough, and P. J. Williams. OPT++: an object oriented toolkit for nonlinear optimization. *ACM Transactions on Mathematical Software*, 33(2), 2007. [15](#)
- [70] A. Narayan and D. Xiu. Stochastic collocation methods on unstructured grids in high dimensions via interpolation. *SIAM J. Scientific Computing*, 34(3), 2012. [36](#)
- [71] L. W. T. Ng and M. S. Eldred. Multifidelity uncertainty quantification using nonintrusive polynomial chaos and stochastic collocation. In *Proceedings of the 14th AIAA Non-Deterministic Approaches Conference*, number AIAA-2012-1852, Honolulu, HI, April 23–26, 2012. [42](#)
- [72] F. Nobile, R. Tempone, and C. G. Webster. A sparse grid stochastic collocation method for partial differential equations with random input data. Technical Report Technical report TRITA-NA 2007:7, Royal Institute of Technology, Stockholm, Sweden, 2007. [32](#)
- [73] F. Nobile, R. Tempone, and C. G. Webster. An anisotropic sparse grid stochastic collocation method for partial differential equations with random input data. *SIAM J. on Num. Anal.*, 46(5):2411–2442, 2008. [32](#)
- [74] J. Nocedal and Wright S. J. *Numerical Optimization*. Springer Series in Operations Research. Springer, New York, 1999. [58](#), [71](#)
- [75] E. O. Omojokun. *Trust Region Algorithms for Optimization with Nonlinear Equality and Inequality Constraints*. PhD thesis, University of Colorado, Boulder, Colorado, 1989. [59](#)
- [76] M.R. Osborne, B. Presnell, and B.A. Turlach. A new approach to variable selection in least squares problems. *IMA Journal of Numerical Analysis*, 20(3):389–403, 2000. [35](#)
- [77] V. M. Pérez, M. S. Eldred, , and J. E. Renaud. Solving the infeasible trust-region problem using approximations. In *Proceedings of the 10th AIAA/ISSMO Multidisciplinary Analysis and Optimization Conference*, Albany, NY, Aug. 30–Sept. 1, 2004. AIAA Paper 2004-4312. [59](#), [61](#)
- [78] V. M. Pérez, J. E. Renaud, and L. T. Watson. An interior-point sequential approximation optimization methodology. *Structural and Multidisciplinary Optimization*, 27(5):360–370, July 2004. [58](#), [59](#), [61](#), [63](#)
- [79] R. Rackwitz. Optimization and risk acceptability based on the Life Quality Index. *Struct. Saf.*, 24:297–331, 2002. [66](#)

- [80] R. Rackwitz and B. Fiessler. Structural reliability under combined random load sequences. *Comput. Struct.*, 9:489–494, 1978. [10](#)
- [81] P. Ranjan, Bingham D., and Michailidis G. Sequential experiment design for contour estimation from complex computer codes. *Technometrics*, 50(4):527–541, 2008. [19](#)
- [82] M. T. Reagan, H. N. Najm, P. P. Pebay, O. M. Knio, and R. G. Ghanem. Quantifying uncertainty in chemical systems modeling. *Int. J. Chem. Kinet.*, 37:368–382, 2005. [38](#)
- [83] G. Rennen. Subset selection from large datasets for kriging modeling. *Structural and Multidisciplinary Optimization*, 38(6):545–569, 2009. [49](#)
- [84] J. F. Rodriguez, J. E. Renaud, and L. T. Watson. Convergence of trust region augmented lagrangian methods using variable fidelity approximation data. *Structural Optimization*, 15:1–7, 1998. [57](#)
- [85] M. Rosenblatt. Remarks on a multivariate transformation. *Annals of Mathematical Statistics*, 23(3):470–472, 1952. [10](#), [30](#)
- [86] J. Sacks, S. B. Schiller, and W. Welch. Design for computer experiments. *Technometrics*, 31:41–47, 1989. [17](#), [18](#)
- [87] I.C. Simpson. Numerical integration over a semi-infinite interval, using the lognormal distribution. *Numerische Mathematik*, 31:71–76, 1978. [22](#)
- [88] S.A. Smolyak. Quadrature and interpolation formulas for tensor products of certain classes of functions. *Dokl. Akad. Nauk SSSR*, 4:240–243, 1963. [32](#)
- [89] A. Stroud. *Approximate Calculation of Multiple Integrals*. Prentice Hall, 1971. [33](#)
- [90] G. Tang, G. Iaccarino, and M. S Eldred. Global sensitivity analysis for stochastic collocation expansion. In *Proceedings of the 51st AIAA/ASME/ASCE/AHS/ASC Structures, Structural Dynamics, and Materials Conference (12th AIAA Non-Deterministic Approaches conference)*, Orlando, FL, April 12–15, 2010. AIAA Paper 2010-XXXX. [39](#)
- [91] M.A. Tatang. *Direct incorporation of uncertainty in chemical and environmental engineering systems*. PhD thesis, MIT, 1995. [34](#)
- [92] Robert Tibshirani. Regression shrinkage and selection via the lasso. *Journal of the Royal Statistical Society, Series B*, 58:267–288, 1996. [35](#)
- [93] J. Tu, K. K. Choi, and Y. H. Park. A new study on reliability-based design optimization. *J. Mech. Design*, 121:557–564, 1999. [11](#)
- [94] A. van der Sluis. Condition numbers and equilibration of matrices. *Numerische Mathematik*, 14(1):14–23, 1969. [54](#)
- [95] G. N. Vanderplaats. *Numerical Optimization Techniques for Engineering Design: With Applications*. McGraw-Hill, New York, 1984. [58](#), [60](#)
- [96] R. W. Walters. Towards stochastic fluid mechanics via polynomial chaos. In *Proceedings of the 41st AIAA Aerospace Sciences Meeting and Exhibit*, number AIAA-2003-0413, Reno, NV, January 6–9, 2003. [34](#)
- [97] G. W. Wasilkowski and H. Woźniakowski. Explicit cost bounds of algorithms for multivariate tensor product problems. *Journal of Complexity*, 11:1–56, 1995. [32](#)
- [98] N. Wiener. The homogeneous chaos. *Amer. J. Math.*, 60:897–936, 1938. [21](#)

- [99] J. A. S. Witteveen and H. Bijl. Modeling arbitrary uncertainties using gram-schmidt polynomial chaos. In *Proceedings of the 44th AIAA Aerospace Sciences Meeting and Exhibit*, number AIAA-2006-0896, Reno, NV, January 9–12 2006. [22](#)
- [100] Y.-T. Wu, Y. Shin, R. Sues, and M. Cesare. Safety-factor based approach for probability-based design optimization. In *Proc. 42nd AIAA/ASME/ASCE/AHS/ASC Structures, Structural Dynamics, and Materials Conference*, number AIAA-2001-1522, Seattle, WA, April 16–19 2001. [67](#), [71](#)
- [101] Y.-T. Wu and P.H. Wirsching. A new algorithm for structural reliability estimation. *J. Eng. Mech., ASCE*, 113:1319–1336, 1987. [10](#)
- [102] B. A. Wujek and J. E. Renaud. New adaptive move-limit management strategy for approximate optimization, part 1. *AIAA Journal*, 36(10):1911–1921, 1998. [59](#)
- [103] B. A. Wujek and J. E. Renaud. New adaptive move-limit management strategy for approximate optimization, part 2. *AIAA Journal*, 36(10):1922–1934, 1998. [59](#)
- [104] D. Xiu. Numerical integration formulas of degree two. *Applied Numerical Mathematics*, 58:1515–1520, 2008. [33](#)
- [105] D. Xiu and J.S. Hesthaven. High-order collocation methods for differential equations with random inputs. *SIAM J. Sci. Comput.*, 27(3):1118–1139 (electronic), 2005. [32](#)
- [106] S. Xu and R. V. Grandhi. Effective two-point function approximation for design optimization. *AIAA J.*, 36(12):2269–2275, 1998. [13](#)
- [107] Hui Zou and Trevor Hastie. Regularization and variable selection via the elastic net. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 67(2):301–320, 2005. [35](#)
- [108] T. Zou, S. Mahadevan, and R. Rebba. Computational efficiency in reliability-based optimization. In *Proceedings of the 9th ASCE Specialty Conference on Probabilistic Mechanics and Structural Reliability*, Albuquerque, NM, July 26–28, 2004. [67](#)
- [109] T. Zou, Z. Mourelatos, S. Mahadevan, and P. Meernik. Reliability analysis of automotive body-door subsystem. *Rel. Eng. and Sys. Safety*, 78:315–324, 2002. [16](#)

DISTRIBUTION:

1 MS 0899 Technical Library, 9536 (electronic copy)

