

SANDIA REPORT

SAND2014-2864

Unlimited Release

Printed March 2014

User Guidelines and Best Practices for CASL VUQ Analysis Using Dakota

Brian M. Adams, Russell W. Hooper, Allison Lewis, Jerry A. McMahan, Jr.,
Ralph C. Smith, Laura P. Swiler, Brian J. Williams

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.



Sandia National Laboratories

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>



SAND2014-2864
Unlimited Release
Printed March 2014

User Guidelines and Best Practices for CASL VUQ Analysis Using Dakota

Brian M. Adams
Laura P. Swiler
Optimization and Uncertainty Quantification Department

Russell W. Hooper
Multiphysics Applications Department

Sandia National Laboratories
P.O. Box 5800
Albuquerque, NM 87185-1318

Allison Lewis
Jerry A. McMahan, Jr.
Ralph C. Smith
Department of Mathematics
North Carolina State University
Box 8205
Raleigh, NC 27695-8205

Brian J. Williams
Statistical Sciences Group (CCS-6)
Los Alamos National Laboratory
Bikini Atoll Road, MS F600
Los Alamos, NM 87545

Abstract

Sandia's Dakota software (available at <http://dakota.sandia.gov>) supports science and engineering transformation through advanced exploration of simulations. Specifically it manages and

analyzes ensembles of simulations to provide broader and deeper perspective for analysts and decision makers. This enables them to enhance understanding of risk, improve products, and assess simulation credibility.

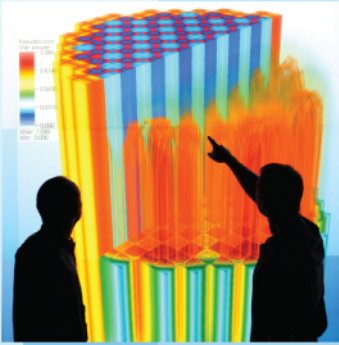
This manual offers Consortium for Advanced Simulation of Light Water Reactors (LWRs) (CASL) partners a guide to conducting Dakota-based VUQ studies for CASL problems. It motivates various classes of Dakota methods and includes examples of their use on representative application problems. On reading, a CASL analyst should understand why and how to apply Dakota to a simulation problem.

This SAND report constitutes the product of CASL milestone L3:VUQ.V&V.P8.01 and is also being released as a CASL unlimited release report with number CASL-U-2014-0038-000.



Power uprates
and plant life extension

L3:VUQ.V&V.P8.01



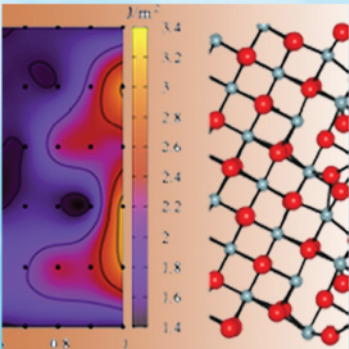
Engineering design
and analysis

User Guidelines and Best Practices for CASL VUQ Analysis Using Dakota

CASL-U-2014-0038-000



Science-enabling
high performance
computing



Fundamental science



Plant operational data

Brian M. Adams¹
Russell W. Hooper¹
Allison Lewis²
Jerry A. McMahan, Jr.²
Ralph C. Smith²
Laura P. Swiler¹
Brian J. Williams³

¹Sandia National Laboratories

²North Carolina State University

³Los Alamos National Laboratory

March 28, 2014



U.S. DEPARTMENT OF
ENERGY

Nuclear Energy

User Guidelines and Best Practices for CASL VUQ Analysis Using Dakota

Brian M. Adams ¹
Russell W. Hooper ²
Allison Lewis ³
Jerry A. McMahan, Jr. ⁴
Ralph C. Smith ⁵
Laura P. Swiler ⁶
Brian J. Williams ⁷

May 6, 2014

¹briadam@sandia.gov

²rhoope@sandia.gov

³allewis2@ncsu.edu

⁴jamcmaha@ncsu.edu

⁵rsmith@ncsu.edu

⁶lpswile@sandia.gov

⁷brianw@lanl.gov

Contents

1	Overview	1
1.1	Manual Contents	2
1.2	Getting Started with Dakota	4
1.3	Acknowledgments	7
2	Application Example Problems	9
2.1	Cantilever beam	9
2.2	General Linear Model Verification Test Suite	11
2.3	COBRA-TF Thermal-Hydraulics Simulation Problem	14
2.3.1	COBRA-TF Simulator Overview	14
2.3.2	COBRA-TF test problem description	15
2.3.3	VUQ Parameters in COBRA-TF Problem 6	17
3	Sensitivity Analysis	19
3.1	Terminology	19
3.1.1	Local Versus Global Sensitivity	19
3.1.2	Sensitivity Metrics	20
3.2	Recommended Methods	23
3.2.1	Centered Parameter Study	24
3.2.2	Multidimensional Parameter Study	26
3.2.3	Global LHS Sampling	31
3.2.4	PSUADE/Morris Method	34
3.3	Summary and Additional Approaches	36
4	Surrogate Models	38
4.1	Polynomial Regression Models	40
4.1.1	Fitting Polynomial Surrogates in Dakota	41
4.2	Kriging and Gaussian Process Models	44
4.2.1	Fitting Kriging Surrogates in Dakota	47
4.3	Summary	48
5	Optimization and Deterministic Calibration	52
5.1	Terminology and Problem Formulations	53
5.1.1	Special Considerations for Calibration	54
5.2	Recommended Methods	56

5.2.1	Gradient-Based Local Methods	56
5.2.2	Derivative-Free Local Methods	60
5.2.3	Derivative-Free Global Methods	65
5.3	Summary and Additional Approaches	68
6	Uncertainty Quantification	70
6.1	Uncertainty Propagation	71
6.1.1	Sampling Methods	71
6.1.2	Stochastic Polynomial Methods	72
6.1.3	Verification	74
6.1.4	Prediction Intervals	75
6.1.5	Uncertainty Propagation: Cantilever Beam Example	75
6.2	Bayesian Model Calibration	83
6.2.1	Direct Implementation of Bayes' Relation	84
6.2.2	Sampling Based Metropolis Algorithms	86
6.2.3	Model Calibration and Surrogate Models	87
6.2.4	Verification	88
6.2.5	Synthetic Data	88
6.2.6	Bayesian Calibration Examples	89
7	COBRA-TF VUQ Studies	103
7.1	Initial Parameter Studies with Two Power Distributions	104
7.2	COBRA-TF Sensitivity Studies	104
7.2.1	Centered Parameter Study	104
7.2.2	Latin hypercube sampling studies	109
7.2.3	Morris Screening	113
7.2.4	Screening to Reduce Parameters	113
7.3	Calibration Studies	116
7.3.1	Deterministic Calibration	117
7.3.2	Surrogate Construction	117
7.3.3	Bayesian Calibration	122
A	General Linear Model Verification Test Suite	129
A.1	Verification Scenarios	130
A.2	Verification Tests	133
B	Procedure for Running COBRA-TF Studies	135

Chapter 1

Overview

Sandia's Dakota software (available at <http://dakota.sandia.gov>) supports science and engineering transformation through advanced exploration of simulations. Specifically it manages and analyzes ensembles of simulations to provide broader and deeper perspective for analysts and decision makers. This enables them to enhance understanding of risk, improve products, and assess simulation credibility.

In its simplest mode, Dakota can automate typical parameter variation studies through a generic interface to a physics-based computational model. This can lend efficiency and rigor to manual parameter perturbation studies already being conducted by analysts. However, Dakota also delivers advanced parametric analysis techniques enabling design exploration, optimization, model calibration, risk analysis, and quantification of margins and uncertainty with such models. It directly supports verification and validation activities. Dakota algorithms enrich complex science and engineering models, enabling an analyst to answer crucial questions of

- **Sensitivity:** Which are the most important input factors or parameters entering the simulation, and how do they influence key outputs?
- **Uncertainty:** What is the uncertainty or variability in simulation output, given uncertainties in input parameters? How safe, reliable, robust, or variable is my system? (Quantification of margins and uncertainty, QMU)
- **Optimization:** What parameter values yield the best performing design or operating condition, given constraints?
- **Calibration:** What models and/or parameters best match experimental data?

In general, Dakota is the Consortium for Advanced Simulation of Light Water Reactors (CASL) delivery vehicle for verification, validation, and uncertainty quantification (VUQ) algorithms. It permits ready application of the VUQ methods described above to simulation codes by CASL researchers, code developers, and application engineers.

More specifically, the CASL VUQ Strategy [26] prescribes the use of Predictive Capability Maturity Model (PCMM) assessments [30]. PCMM is an expert elicitation tool designed to characterize and communicate completeness of the approaches used for computational model definition, verification, validation, and uncertainty quantification associated with an intended application. Exercising

a computational model with the methods in Dakota will yield, in part, evidence for a predictive capability maturity model (PCMM) assessment. Table 1.1 summarizes some key predictive maturity related activities (see details in [26]), with examples of how Dakota fits in.

Table 1.1: Summary of Dakota relevance for PCMM-related activities.

VUQ/PCMM Activity	Dakota relevance
select quantities of interest (QOIs)	(limited)
software quality assurance (SQA)	(limited)
code verification	conduct parameter studies as a function of mesh quality; calculate convergence rate
solution verification	conduct parameter studies over solver parameters or mesh quality; calculate convergence rate
validation	run ensemble of simulations and make (potentially uncertainty-aware) comparisons with experimental data, assess model form uncertainty
sensitivity	conduct global sensitivity analysis to rank or screen parameters; supports quantified parameter ranking table (QPRT)
uncertainty quantification	compute uncertainty in QOIs for risk-informed decision making
calibration	tune or refine models for use in particular scenarios

This manual offers CASL partners a guide to conducting Dakota-based VUQ studies for CASL problems. It motivates various classes of Dakota methods and includes examples of their use on representative application problems. On reading, a CASL analyst should understand why and how to apply Dakota to a simulation problem.

1.1 Manual Contents

This user’s guide emphasizes best practice and highlights a few key approaches to solving the VUQ analysis problems of greatest interest to CASL today. The remainder of this chapter summarizes high level steps to getting started using Dakota. Chapter 2 describes simple, but physically meaningful, application problems that will be used to demonstrate each of the Dakota algorithmic approaches described in subsequent chapters. Among the examples is a CASL-relevant COBRA-TF thermal-hydraulic simulator problem.

Chapters 3 through 6 tour four of Dakota’s major algorithmic capabilities. In each chapter you will find high-level analysis goals and terminology, with references to more detailed descriptions and theory; guidance on how to choose from the available Dakota approaches and assess whether they are working; and application examples, including Dakota input, Dakota output, and post-processing/interpretation. The major VUQ activities addressed by this manual are:

- **Parameter Studies and Sensitivity Analysis:** Dakota parameter studies automate typical parameter variation studies such as running the model at a tensor grid of parameter values or varying them 1%, 5%, 10% from a nominal value. Sensitivity analysis determines model parameters most influential on quantities of interest (responses). This can be used to rank the influence of parameters, as in a Quantified Parameter Ranking Table (QPRT) [26], or screen/down-select to a tractable number of free parameters for follow-on analyses. See Chapter 3, Sensitivity Analysis for an overview of parameter studies, global sensitivity analysis methods and metrics, and a demonstration of using Dakota to perform parameter ranking.
- **Surrogate Models:** Any of the Dakota studies described can be conducted directly on a computational model or with surrogate model indirection. Surrogate models are inexpensive approximate models that are intended to capture the salient features of an expensive high-fidelity model. In this manual and the Dakota context, surrogate models are not based on simplifying physical assumptions. Rather they are response surface models constructed automatically by Dakota based on empirical samples of the true simulation's input/output behavior. For example, in CASL one might run costly CFD simulations at a set of design points in a parameter space and then have Dakota build an algebraic Kriging model on the QOI data for use in optimization. Chapter 4 has an overview of the most commonly used surrogate models, which can smooth noisy model responses, or reduce computational cost. On reading it, you will be able to create Dakota studies that automatically run a computational model, generate a response surface model, and evaluate it in the context of another Dakota study.
- **Calibration:** Dakota provides capabilities for automatically tuning model parameters to best match experimental (or high-fidelity model) data. This process is also known as parameter estimation, calibration, data assimilation, or model inversion to update knowledge of parameter values based on additional data. A CASL example would be tuning crud chemistry reaction rates to match experimental data. Approaches yielding single point estimates of parameters are described in Chapter 5, Optimization and Deterministic Calibration, while Bayesian methods resulting in a probability distribution for the unknown parameters are covered in Section 6.2, Bayesian Model Calibration.
- **Design Optimization:** Adjusting model parameters to meet desired performance criteria while satisfying other constraints. For example, determine optimal shape to minimize vibration, or design mixing vanes to minimize crud formation. Chapter 5, Optimization and Deterministic Calibration will help you choose from among Dakota optimization methods based on problem characteristics and your specific optimization goals.
- **Uncertainty Quantification (UQ):** Model predictions with quantified uncertainty support validation and follow-on decision making. UQ methods accept characterizations of input parameter uncertainty and run the computational model to compute the resulting uncertainties on response quantities of interest. UQ methods, which yield statistics on QOIs (mean, standard deviation, distribution, range), are described in Chapter 6, Uncertainty Quantification.

This guide selectively focuses on two to three ways to execute each type of VUQ analysis depending on goals and problem characteristics. The approaches included have worked well in practice on a broad range of problems in computational science and engineering. When challenges

are encountered, many alternative and advanced methods that may perform better are available in Dakota.

The manual concludes with a realistic thermal-hydraulics example from a CASL “Progression Problem” in Chapter 7, COBRA-TF VUQ Studies. The example demonstrates a VUQ activity flow from initial parameter studies, through sensitivity analysis for parameter screening, to calibration on a reduced parameter set. Construction of the surrogate utilized in model calibration as a substitute for more costly direct COBRA-TF calculations is also illustrated.

Additional Resources

This user’s guide is not an exhaustive guide to Dakota’s capabilities. It is a high-level supplement to other Dakota and VUQ resources. Users reaching its extent should consult:

- The Dakota User’s Manual [1]: a more complete summary of Dakota capabilities from getting started through advanced methods (<http://dakota.sandia.gov/docs/dakota/5.4/Users-5.4.pdf>);
- The Dakota Reference Manual [2]: extensive guidance on valid keywords to use in a Dakota input file to specify a Dakota study (<http://dakota.sandia.gov/docs/dakota/5.4/html-ref/index.html>);
- The directories `dakota/examples` and `dakota/test` included with Dakota distributions which contain examples of input files referenced in the documentation and many more (also available at <https://software.sandia.gov/trac/dakota/browser/tags/5.4/examples> and <https://software.sandia.gov/trac/dakota/browser/tags/5.4/test>);
- The Dakota Theory Manual [3] (<http://dakota.sandia.gov/docs/dakota/5.4/Theory-5.4.pdf>) and research publications available at <http://dakota.sandia.gov/publications.html>: detailed background on algorithmic approaches developed directly in Dakota to tackle challenging science and engineering analyses; and
- Publications referenced throughout all the above.

This document refers to Dakota 5.4 and its documentation; newer versions may be available on the Dakota website. This guide for Dakota usage in CASL aims to be generic and thus does not supplant any domain-specific best practices or guidance for performing VUQ-related studies.

1.2 Getting Started with Dakota

The remainder of this manual largely focuses on selecting and applying Dakota methods and understanding the results. This section surveys at a higher level some prerequisites for using Dakota, with references to additional resources for help.

Know Why to Use Dakota

Understanding your simulation’s characteristics, your VUQ analysis goals, and Dakota’s relevance in achieving them are critical first steps. These likely seem obvious, but are crucial in order to select from the many available methods in Dakota. This guide aims to address this background by offering

a high-level introduction to some key analysis methods, their application, and benefits. Other resources to understanding Dakota’s applicability include training materials, publicity materials, and publications on the Dakota website <http://dakota.sandia.gov>, as well as the Dakota User’s Manual [1].

Access the Software and Other Resources

Dakota: Dakota is available to CASL partners as part of the Virtual Environment for Reactor Applications (VERA). It can be checked out via git clone from `casl-dev.ornl.gov:/git-root/Dakota`. Dakota is also distributed with VERA releases under `Trilinos/packages/TriKota/Dakota`. Examples of Dakota applied to CASL problems are evolving. These are archived in milestone reports and protected in the VERA software in the VUQDemos suite, available via git clone from `casl-dev.ornl.gov:/git-root/VUQDemos`. Dakota also has a public download site <http://dakota.sandia.gov/download.html>, which may be useful for CASL partners without ready access to the VERA development environment at Oak Ridge National Laboratory (ORNL).

This Manual and Examples: Examples from this manual, including input, output, auxiliary data files, and scripts are available from the CASL Git repository at `casl-dev.ornl.gov:/git-root/VUQDemos/CaslDakotaManual`. The L^AT_EX source and images are also included.

Help Resources: For help beyond this manual and the documents referenced herein, see the Dakota website: <http://dakota.sandia.gov>. It includes software downloads, documentation, publications, and training materials. It also has guidance on seeking general help with Dakota via the `dakota-users` mailing list (`dakota-users@software.sandia.gov`). CASL-specific issues should be directed to `casl-vuq@casl.gov`.

How to Interact with Dakota

An overall Dakota analysis process is depicted in Figure 1.1. The specification of the VUQ analysis problem to Dakota is given in the text-based Dakota input file, including the method (algorithm) which dictates how Dakota generates parameter sets at which to run the user’s simulation. As Dakota runs, it will iteratively evaluate the simulation at these parameter sets by running a user-provided analysis driver and collecting corresponding quantities of interest output by the simulation workflow. When complete, the Dakota executable will produce console text output and tabular data with VUQ results for subsequent analysis.

Interface Dakota to the Simulation

Dakota requires an analysis driver to communicate with the computational model. The contract for this Dakota/simulation interface is straightforward: it must be an automated workflow that accepts Dakota parameters as input from a text file, runs the simulation in a batch/non-interactive mode, and produces responses (quantities of interest derived from simulation output) in a text file for consumption by Dakota.

As Dakota runs, it will determine values of parameters for which response data is needed. When ready to evaluate the simulation at such a parameter set, Dakota will write a “parameters file” with the values of the variables. Dakota will then invoke the specified analysis driver, represented by the dashed blue box in Figure 1.1. This analysis driver must implement the automated process that

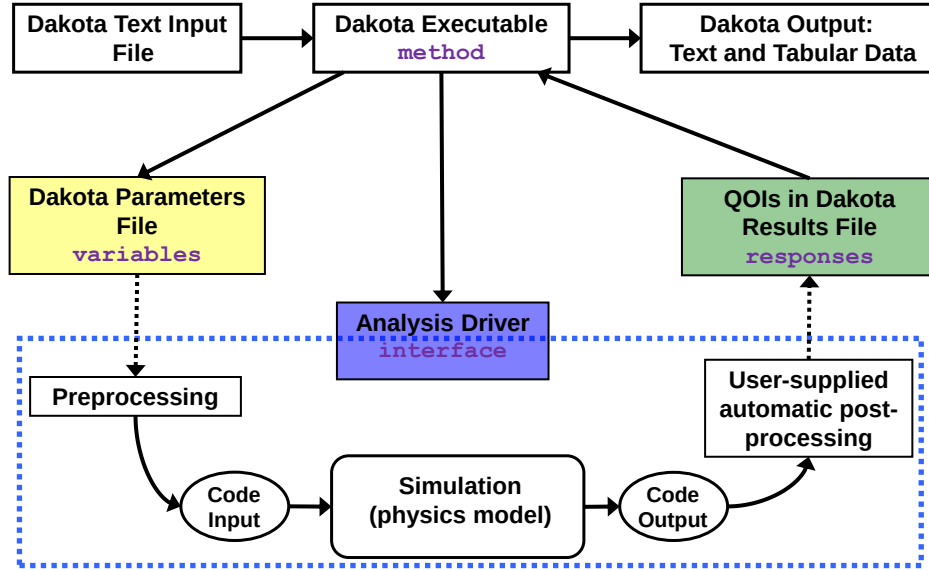


Figure 1.1: Components of a Dakota study, including Dakota input and output, and interface to a computational model (simulation).

takes a Dakota parameters file as input and produces a Dakota results file as output. Typically this driver is a script which includes preprocessing, running the code, and postprocessing to extract QOIs from simulation output. When the driver completes, Dakota requires the “results file” to contain the quantities of interest resulting from running the simulation at the specified parameter values.

Additional resources for creating and running a Dakota/simulation workflow include:

- High level guidance in “User Supplied Simulation Code Examples” in the Dakota Tutorial in the Dakota User’s Manual [1], with more details in the “Advanced Simulation Code Interfaces” chapter, together with example code in `dakota/examples/script_interfaces`.
- Considerations for having Dakota manage concurrent simulation runs in parallel, as this is a common need. Managing parallel concurrency locally, within queue, and out of queue, including batch submission and later retrieval are addressed by the Dakota User’s Manual, “Application Parallelism Use Cases”, together with examples in `dakota/examples/parallelism`.
- CASL-specific Dakota workflows with COBRA-TF and Insilico, demonstrated in the above referenced VUQDemos.

Understand Dakota Input Files

Once an interface is constructed between Dakota and the simulation, one may readily apply any Dakota method by simply changing the input file. Dakota input files are simple plain text files with six categories of information that can appear (some are optional and some may appear multiple times in advanced studies). Four of these are indicated notionally in Figure 1.1:

- **strategy** (not depicted): overall control of Dakota methods and tabular output data.

- **method:** specifies the iterative analysis method being run on the model, for example Latin hypercube sampling or gradient-based parameter estimation.
- **variables:** characterization of the model parameters Dakota is varying in the study, such as lognormal uncertain or continuous design, together with supplementary fixed (state) parameters.
- **responses:** quantities of interest returned to Dakota for analysis.
- **interface:** the simulation workflow mapping variables to responses in an automated way; typically a script that orchestrates this workflow.
- **model** (not depicted): a container encapsulating a set of variables, interface, and responses for presentation to a method; useful for specifying that a surrogate model should be automatically constructed to serve as a proxy for an expensive computational model.

This guide shows and explains a number of examples of input files which configure Dakota to conduct various kinds of iterative analyses. They can be taken verbatim from the text to conduct Dakota studies and will also be available from CASL records systems when this manual is published. Additional examples are available in the various Dakota manuals and with the Dakota software itself. Specific guidance on individual Dakota keywords is available in the Dakota Reference Manual [2], which can help when trying to determine how to configure Dakota for a new kind of study.

Understand Dakota Results

This document offers an introduction to Dakota output, including log file output and tabular data to understand the results of studies. Often this data must be interpreted or post-processed with external tools to be useful in a decision making context. Examples are included in this manual to demonstrate this.

1.3 Acknowledgments

This CASL/Dakota manual borrows heavily from the Dakota 5.4 User's Manual [1]. We heartily thank the authors of the most recent version of that document: 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, and Dena M. Vigil.

Dakota's use of the Quantification of Uncertainty for Estimation, Simulation, and Optimization (QUESO) library is facilitated through close interaction with its developers at the University of Texas at Austin. Ernesto E. Prudencio, Nicholas Malaya, and Damon McDougall have kindly provided examples and test problems implemented in QUESO.

We also appreciate the efficient implementation of code changes to COBRA-TF made by Noel K. Belcourt to expose code parameters to Dakota needed to drive the various COBRA-TF parameter studies described in this manual.

We appreciate the valuable feedback provided by several reviewers of this manual, including Patty Hough, Vince Mousseau, Rod Schmidt, and Dena Vigil.

This research was supported by the Consortium for Advanced Simulation of Light Water Reactors (<http://www.casl.gov>), an Energy Innovation Hub (<http://www.energy.gov/hubs>) for

Modeling and Simulation of Nuclear Reactors under U.S. Department of Energy Contract No. DE-AC05-00OR22725. Sandia National Laboratories (SNL), North Carolina State University (NCSU), and Los Alamos National Laboratory (LANL) are core CASL partners.

Chapter 2

Application Example Problems

This chapter describes representative, though simplified, application problems that will be used to demonstrate various Dakota approaches. Each of the three application examples has strengths to help bridge abstract Dakota concepts and guidance throughout the remainder of this manual to concrete practice:

- **Cantilever beam:** A simple static mechanics analysis where prescribed geometry, material properties, and loads on a cantilevered beam map to quantities of interest such as weight, displacement, and stress. The physical meaning is intuitive, the physics equations have a simple algebraic form, and a simulator for it is included with Dakota, along with several example input files. Each Dakota analysis technique in Chapters 3 through 6 includes a worked example using the cantilever beam problem.
- **General linear model:** A model with closed algebraic form specifically designed to support algorithm and code verification, i.e., to verify that VUQ algorithms are working as expected. This example consists of a linear mapping from model parameters to responses, with an additive noise term. It is used to verify the performance of Bayesian calibration methods in Section 6.2.
- **COBRA-TF thermal-hydraulics:** A coupled physics, single assembly reactor model problem simulated with CASL's COBRA-TF thermal-hydraulics code. The most physically realistic example, this problem demonstrates VUQ infrastructure for varying model form and other parameters, running a computational model, and distilling quantities of interest from code output. This example is the focus of Chapter 7, where it is used to demonstrate parameter screening, calibration, and surrogate construction.

2.1 Cantilever beam

The cantilever beam example problem is adapted from the reliability-based design optimization literature [36], [42]. The uniform cantilever beam is shown in Figure 2.1, with a left anchor and a fixed length $L = 100\text{in}$. The beam width and thickness are parameterized by w and t , respectively. The free end of the beam is subject to horizontal load X and vertical load Y .

Given Young's elastic modulus E , the simplified algebraic physics equations used to model the

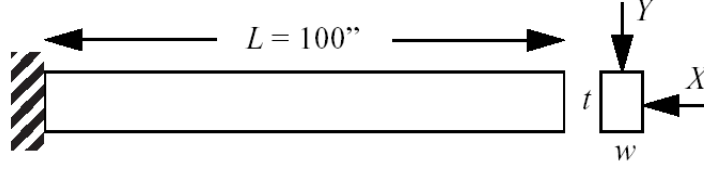


Figure 2.1: Cantilever beam test problem.

beam area A (a stand-in for weight W), stress S , and displacement D are:

$$\begin{aligned} W &\propto A = wt \\ S &= \frac{600}{wt^2}Y + \frac{600}{w^2t}X \\ D &= \frac{4L^3}{Ewt} \sqrt{\left(\frac{Y}{t^2}\right)^2 + \left(\frac{X}{w^2}\right)^2} \end{aligned} \quad (2.1)$$

Given a specified maximum displacement D_0 and yield stress R , the quantities of interest for the Dakota study (Dakota “responses”) are defined to be:

$$\begin{aligned} \text{area} &= wt \\ \text{stress} &= S - R = \frac{600}{wt^2}Y + \frac{600}{w^2t}X - R \\ \text{displacement} &= D - D_0 = \frac{4L^3}{Ewt} \sqrt{\left(\frac{Y}{t^2}\right)^2 + \left(\frac{X}{w^2}\right)^2} - D_0, \end{aligned} \quad (2.2)$$

where **area** is used as a surrogate for weight, and **stress** and **displacement** are defined as differences with respect to the prescribed values of stress and displacement. The parameters in the problem are summarized with their nominal values in Table 2.1. Parameters w and t have simple bounds, while R, E, X , and Y may be considered as random variables and modeled with normal distributions $\mathcal{N}(\mu, \sigma^2)$ when conducting uncertainty studies.

Table 2.1: Parameters for the cantilever beam problem.

parameter	description	nominal value	range/distribution
w	width	2.5	[1.0, 4.0]
t	thickness	2.5	[1.0, 4.0]
R	yield stress	40000	$\mathcal{N}(4.0E5, 4.0E6)$
E	Young’s modulus	2.9E7	$\mathcal{N}(2.9E7, 2.1025E12)$
X	horizontal load	500	$\mathcal{N}(500, 1.0E4)$
Y	vertical load	1000	$\mathcal{N}(1.0E3, 1.0E4)$
L	beam length	100	-
D_0	maximum displacement	2.2535	-

The built-in Dakota analysis driver `mod_cantilever` computes the outputs **area**, **stress**, and **displacement** given specified inputs w, t, R, E, X , and Y . These cantilever input/output mappings

will be utilized throughout this manual to illustrate the application of core CASL VUQ technologies with Dakota. In Section 3.2, the sensitivity of the quantities of interest with respect to the input parameters is assessed to rank their importance and exercise the model.

Then, in Sections 4.1.1 and 4.2.1, response surface models (surrogates) for the parameter to response mapping are generated based on a small number of sampled runs of cantilever. This emulates the practical process one must use when models are costly. In Section 5.2, a deterministic design optimization problem is solved to design the beam geometry. The goal is to minimize the weight (or equivalently, the cross-sectional area) of the beam subject to a displacement constraint and a stress constraint. The parameters R , E , X , and Y are fixed at their nominal values and the deterministic design problem is given by

$$\begin{aligned}
& \text{minimize} && \text{area} = wt \\
& \text{subject to} && \text{stress} = S - R \leq 0 \\
& && \text{displacement} = D - D_0 \leq 0 \\
& && 1.0 \leq w \leq 4.0 \\
& && 1.0 \leq t \leq 4.0
\end{aligned} \tag{2.3}$$

In Section 5.2.1, the cantilever beam is calibrated to synthetic experimental data for area, stress, and displacement, to find the values of w , t , and E yielding best agreement with the data. Finally, when considered for uncertainty quantification in Chapter 6, the design variables are fixed at their nominal or optimal values, and the Dakota study is conducted over the normally-distributed uncertain parameters R , E , X , and Y . This yields estimates of the mean, standard deviation, and overall distribution of the quantities of interest.

2.2 General Linear Model Verification Test Suite

Whereas most CASL codes exhibit a nonlinear input-output relation, linearly parameterized problems serve an important role for algorithm and code verification. These uses include the following:

- They provide a hierarchy of models, which can be used to test the convergence of Bayesian model calibration algorithms through comparison with analytic solutions.
- They provide a regime to test of the accuracy of uncertainty propagation algorithms since one can employ analytic relations between input and output densities.
- They facilitate the testing of algorithms for heavy-tailed distributions.
- They provide a framework for analytically testing algorithms to construct Sobol global sensitivity indices (described in Section 3.1.2).

The family of linear models described in this section is used to verify the performance of Bayesian calibration methods in Section 6.2. A high-level overview of the problem appears here, with additional details in Appendix A. A simulator implementing this problem is available on request from the authors.

We employ the linear regression model

$$Y = G\beta + \varepsilon(\lambda, \phi). \tag{2.4}$$

as a test problem for verifying the Bayesian calibration capabilities in Dakota, which currently include Quantification of Uncertainty for Estimation, Simulation, and Optimization (QUESO) and Differential Evolution with Self-Adaptive Randomized Subspace Sampling (via the DREAM software package). In this model Y is the N -dimensional vector of noisy observations of the quantity of interest. The $N \times N_\beta$ matrix G is known and the N_β components of the vector β are unknown regression parameters to be estimated. The N -dimensional random variable ε represents the random measurement noise in the observations. The measurement noise is normally distributed with mean zero and $N \times N$ covariance matrix $R(\phi)/\lambda$. The precision (inverse variance) λ of the ε process is a positive scalar and the permissible values for the correlation structural parameter ϕ depend on the type of correlation being considered as discussed below.

The data for the calibration problem is generated using (2.4) after selecting values for β , λ , and ϕ which are considered the true parameter values (designated β_0 , λ_0 , and ϕ_0). Table 2.2 summarizes three separate cases distinguished by the choice of parameters to be calibrated (or equivalently, by the set of true values considered to be known).

Case	Calibrated	Known
1	β	λ_0, ϕ_0
2	β, λ	ϕ_0
3	β, λ, ϕ	-

Table 2.2: Cases considered for the general linear model verification test suite.

In Case 1, we estimate the regression parameters β assuming the statistics of the measurement noise are perfectly characterized. Case 2 removes explicit knowledge of the error ε precision λ . In Case 3, we assume further that the correlation structure of the measurement noise is only qualitatively known and estimate its parameters ϕ .

The likelihood function used in calibration is proportional to

$$\frac{\lambda^{N/2}}{\det(R(\phi))^{1/2}} \exp \left[-\frac{\lambda}{2} (y - G\beta)^T R^{-1}(\phi) (y - G\beta) \right].$$

For cases where the true value of λ is known, $\lambda = \lambda_0$ is used for computing the likelihood. Similarly, $\phi = \phi_0$ in likelihood calculations when ϕ is known.

Each of the three cases has two subcases defined by whether an informative or noninformative prior is specified for β . The informative prior weights the regression parameter space to indicate a belief that the true parameters are more likely to lie within certain subsets of the parameter domain, while the noninformative prior is agnostic with respect to the location of the regression parameters.

For Case 1, the informative prior for β is specified via a Gaussian random variable having N_β -dimensional mean vector μ and diagonal $N_\beta \times N_\beta$ covariance matrix

$$\frac{1}{\lambda} \begin{bmatrix} \frac{q+r_1}{qr_1} & 0 & \dots & 0 \\ 0 & \frac{q+r_2}{qr_2} & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \frac{q+r_{N_\beta}}{qr_{N_\beta}} \end{bmatrix},$$

where q is a small, positive parameter and $r_1, \dots, r_{N_\beta}$ are positive parameters.

In Case 2, the informative prior for β is the same as in Case 1, and λ_0 is replaced by a random variable λ having prior density proportional to $1/\lambda$. This prior specification is invariant to bijective transformation, i.e. it does not depend on how scale is represented in the error process ε . The prior specification for (β, λ) in Case 3 carries over from Case 2. In Case 3, fixed ϕ_0 is replaced by random ϕ , and the prior distribution for ϕ is specified below.

For Case 1, the noninformative prior for β specifies a uniform density $\pi(\beta) \propto 1$. Case 2 employs the Jeffreys noninformative prior $\pi(\beta, \lambda) \propto 1/\lambda$, which is invariant to bijective transformations of the location and scale parameters. Case 3 also employs the Jeffreys prior for β and λ , while the prior distribution for ϕ is specified below.

Under the assumption of independent and identically distributed errors ε_i , the sample means of the parameters being calibrated converge to their corresponding true values as the number of measurements, N , is increased. Furthermore, for any N and errors ε_i sampled from a mean-zero process having arbitrary covariance structure, the distribution of parameters sampled from QUESO converges to a probability distribution known analytically for (β, λ) , and numerically for ϕ , as the number of QUESO samples is increased. Appendix A provides a more general and rigorous specification of prior distributions for the three cases as well as analytical results for calibrated parameter distributions.

Four types of correlation structure are considered. The types and corresponding domain of ϕ follow:

1. No correlation, no ϕ dependence.
2. Equal correlation, $0 < \phi < 1$.
3. Order 1 autoregressive correlation, $-1 < \phi < 1$.
4. Gaussian spatial correlation, $\phi = (\phi_1, \dots, \phi_M)$, $\phi_i \geq 0$.

For each correlation type, a correlation function and the resulting correlation matrix $R(\phi)$ is provided in Appendix A. The no correlation case indicates no correlation between output measurements and no dependence on ϕ . Equal correlation specifies that all output measurements are equally correlated with ϕ determining the amount of correlation. For order 1 autoregressive correlation, any two output measurements y_i, y_j become less correlated the further apart they are in index (i.e., as $|i - j|$ increases). For Gaussian spatial correlation, the correlation between two output measurements depends on the spatial locations of the corresponding input variables specifying each measurement, with ϕ a vector that adjusts this spatial dependence individually for each input dimension.

The prior distribution for ϕ is taken to be uniform on the allowable domain for ϕ in the equal and autoregressive correlation cases. For Gaussian spatial correlation, assume the M inputs x are restricted by the bounds $\ell_i \leq x_i \leq u_i$ for $i = 1, \dots, M$. The prior distribution $\pi(\phi)$ for ϕ is given by

$$\pi(\phi) \propto \prod_{i=1}^M [\rho_i(\phi_i)]^{a_\rho} [1 - \rho_i(\phi_i)]^{b_\rho - 1} \chi_{[0, \infty)}(\phi_i),$$

where $\rho_i(\phi) = \exp[-\phi(u_i - \ell_i)^2/4]$ for $i = 1, \dots, M$, $(a_\rho, b_\rho) = (1, 0.1)$, and $\chi_{[0, \infty)}(\phi)$ is the characteristic function taking value 1 for $\phi \geq 0$ and 0 for $\phi < 0$.

2.3 COBRA-TF Thermal-Hydraulics Simulation Problem

This section provides an overview of COBRA-TF and a particular thermal-hydraulics simulation problem, CASL VERA Progression Problem 6. An end-to-end demonstration of Dakota methods using this COBRA-TF model is in Chapter 7, COBRA-TF VUQ Studies. The full Dakota/COBRA-TF example is available in the CASL software repositories. Details on accessing it are provided in Appendix B.

2.3.1 COBRA-TF Simulator Overview

COBRA-TF is a thermal-hydraulic (T/H) simulation code designed for light water reactor (LWR) analysis) [5]. COBRA-TF has a long lineage back to the original COBRA computer code developed in 1980 by Pacific Northwest Laboratory, under sponsorship of the Nuclear Regulatory Commission (NRC). The original COBRA began as a thermal-hydraulic rod-bundle analysis code, but subsequent versions have updated and expanded over the past several decades to cover almost all steady-state and transient analyses of both pressurized water reactors (PWRs) and boiling water reactors (BWRs). COBRA-TF is currently developed and maintained by the Reactor Dynamics and Fuel Management Group (RDFMG) at the Pennsylvania State University (PSU). Additional information can be found at the RDFMG website, <http://www.mne.psu.edu/RDFMG/index.html>.

COBRA-TF includes a wide range of thermal-hydraulic models important to LWR safety analysis including flow regime dependent two-phase wall heat transfer, inter-phase heat transfer and drag, droplet breakup, and quench-front tracking. COBRA-TF also includes several internal models to help facilitate the simulation of realistic fuel assemblies. These models include spacer grid models, a fuel rod conduction model, and built-in material properties for both the structural materials and the coolant (i.e., steam tables).

COBRA-TF uses a two-fluid, three-field representation of the two-phase flow. The equations and fields solved are:

- Continuous vapor (mass, momentum and energy)
- Continuous liquid (mass, momentum and energy)
- Entrained liquid drops (mass and momentum)
- Non-condensable gas mixture (mass)

Reasons for selecting COBRA-TF as the primary T/H solver in the VERA core simulator include: reasonable run-times compared to CFD (although CFD will be available as an option), the fact that it is being actively developed and supported by PSU, ability to support future applications of VERA such as transient safety analysis and BWR and SMR applications.

2.3.2 COBRA-TF test problem description

The thermal-hydraulics application example problem used in this manual is a coupled single assembly problem known in CASL as Progression Problem 6 [31]. It simulates a single PWR assembly based on the dimensions and state conditions of Watts Bar Unit 1 Cycle 1. The dimensions for the assembly are identical to AMA Progression Benchmarks “Problem 3” and “Problem 6”. Problems 3 and 6 are identical, except that Problem 3 is at Hot Zero Power (HZIP) and has no T/H feedback, and Problem 6 is at Hot Full Power (HFP) and includes T/H feedback. The test case was run at a boron concentration of 1300 ppm and a 100% power level.

The assembly is a standard 17x17 Westinghouse fuel design with uniform fuel enrichment. There are no axial blankets or enrichment zones. The assembly has 264 fuel rods, 24 guide tubes, and a single instrument tube in the center. There are no control rods or removable burnable absorber assemblies in this problem. The primary geometry specifications of the fuel rod and guide tube materials are given in Figure 2.2 and Table 2.3. The geometry specification for the assembly is given in Figure 2.3 and Table 2.4. The thermal-hydraulic specifications for this problem are shown in Table 2.5.

The COBRA-TF results in this CASL/Dakota manual use a simplified and more efficient adaptation of Progression Problem 6. This study involves only the thermal hydraulic component of Progression Problem 6, holding constant the power supplied by the neutronics component in the full problem. In practice, the neutronics component in the full problem has proved to be at least an order of magnitude more computationally expensive than the thermal hydraulics component. This adapted problem allows relatively rapid and representative sensitivity studies to be performed while expediting testing and refinement of the CASL VUQ software toolset that drives these studies.

Table 2.3: COBRA-TF Problem 6 fuel rod and guide tube descriptions.

Parameter	Value	Units
Fuel Pellet Radius	0.4096	cm
Fuel Rod Clad Inner Radius	0.418	cm
Fuel Rod Clad Outer Radius	0.475	cm
Guide Tube Inner Radius	0.561	cm
Guide Tube Outer Radius	0.602	cm
Instrument Tube Inner Radius	0.559	cm
Instrument Tube Outer Radius	0.605	cm
Outside Rod Height	385.10	cm
Fuel Stack Height (active fuel)	365.76	cm
Plenum Height	16.00	cm
End Plug Heights (x2)	1.67	cm
Pellet Material	UO ₂	
Clad / Caps / Guide Tube Material	Zircaloy-4	

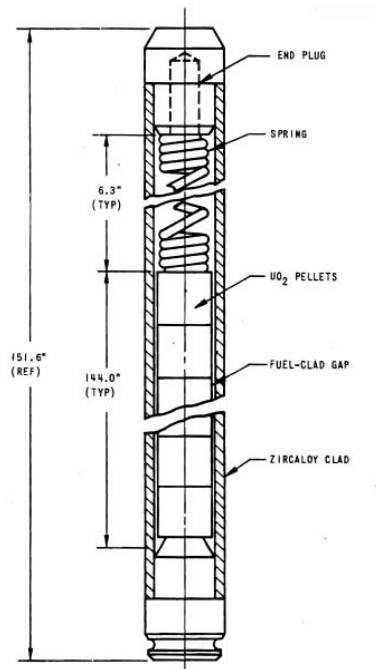


Figure 2.2: COBRA-TF Problem 6 fuel rod diagram.

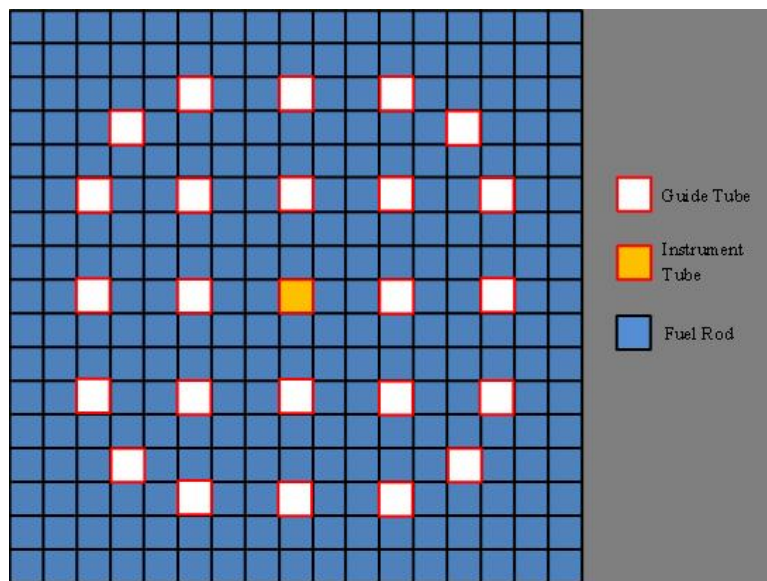


Figure 2.3: COBRA-TF Problem 6 assembly layout showing guide tubes and instrument tube placement.

Table 2.4: COBRA-TF Problem 6 assembly specification.

Parameter	Value	Units
Rod Pitch	1.26	cm
Assembly Pitch	21.5	cm
Inter-Assembly Half Gaps	0.04	cm
Geometry	17x17	
Number of Fuel Rods	264	
Number of Guide Tubes	24	
Number of Instrument Tubes	1	

Table 2.5: COBRA-TF Problem 6 nominal thermal-hydraulic conditions.

Parameter	Value	Units
Inlet Temperature	559	degrees F
System Pressure	2250	psia
Rated Flow (100% flow)	0.6824	Mlb/hr
Rated Power (100% power)	17.67	MWt

2.3.3 VUQ Parameters in COBRA-TF Problem 6

At present, CASL VUQ workflows support COBRA-TF simulation parameter variation via two mechanisms. The first allows Dakota to seamlessly integrate with the VERA Common Input tool suite to perturb any parameters exposed to a user. The second path targets specific code parameters in the thermal hydraulics code that represent all physical phenomena modeled with closure laws in COBRA-TF. These “VUQ parameters” are not exposed to a normal user but are instead exposed to Dakota using an auxiliary input file. The principle is that most analyst users should only perturb input data appearing in the VERA text input file, while advanced VUQ users may need to perturb more advanced parameters such as closure laws. The COBRA-TF studies presented in Chapter 7 use the second mode of parameter variation, i.e. perturbing code parameters via the auxiliary file.

For each parameter, Dakota is able to apply perturbations representing combined shift and scaling, e.g. for an arbitrary parameter p , Dakota can specify values for k_p and ka_p which are used as follows:

$$\bar{p} = k_p * p + ka_p \quad (2.5)$$

The relevant parameters identified by Noel Belcourt and the COBRA-TF code team along with brief descriptions taken from [32] are summarized in Table 2.6. The entries containing “(??)” in their description were not documented in [32] but instead were inferred from the COBRA-TF source code.

Table 2.6: Relevant COBRA-TF thermal-hydraulic code parameters identified by PIRT study.

cd	Pressure loss coefficient of spacer in sub-channel
cdfb	Pressure loss coefficient for sub-channel flow blockage (??)
cond	Thermal conductivity of radial heat transfer
eta	Fraction of vapor generation rate coming from the entrained liquid field
gama	New time vapor generation rate in sub-channel
ql*	Heat transfer rate to liquid in sub-channel
qliht	Heat transfer due to drop impact (??)
qradd	Radiative heat transfer rate from wall to entrained liquid
qradv	Radiative heat transfer rate from wall to vapor
qv*	Heat transfer rate to liquid in sub-channel
qvapl	Incremental heat transferred from grid to vapor (??)
rodqq	Externally supplied heat rate of current rod at current time step (axially averaged)
sdent	Deposition mass flow rate in sub-channel
sent	Entrainment mass flow rate in sub-channel
sphts	Specific heat of radial heat transfer
tmasg	Loss of mass of non-condensable gas in local axial fluid continuity cell due to mixing and void drift to radially adjacent fluid cells
tmasl	Loss of mass of continuous liquid in local axial fluid continuity cell due to mixing and void drift to radially adjacent fluid cells
tmasv	Loss of mass of vapor in local axial fluid continuity cell due to mixing and void drift to radially adjacent fluid cells
tmome	Loss of momentum of droplets in sub-channel due to mixing and void drift to radially adjacent fluid cells
tmoml	Loss of momentum of continuous liquid in sub-channel due to mixing and void drift to radially adjacent fluid cells
tmomv	Loss of momentum of vapor in sub-channel due to mixing and void drift to radially adjacent fluid cells
tnrgl	Loss of enthalpy of liquid in local axial fluid continuity due to mixing and void drift to radially adjacent fluid cells
tnrgv	Loss of enthalpy of vapor in local axial fluid continuity due to mixing and void drift to radially adjacent fluid cells
wkr	Lateral gap pressure loss coefficient
xk	Vertical interfacial drag coefficient between the continuous liquid and vapor phases
xkes	Sink interfacial drag coefficient between the liquid and vapor phases
xkge	Vertical interfacial drag coefficient between the entrained liquid and vapor phases
xkl	Transverse interfacial drag coefficient between the continuous liquid and vapor phases
xkle	Transverse interfacial drag coefficient between the entrained liquid and vapor phases
xkvl	Sink interfacial drag coefficient between the continuous liquid and vapor phases
xkwew	Transverse entrained liquid form loss coefficient
xkwlw	Transverse liquid wall drag coefficient
xkwlx	Vertical liquid wall drag coefficient
xkwvw	Transverse vapor wall drag coefficient
xkwvx	Vertical vapor wall drag coefficient

Chapter 3

Sensitivity Analysis

Broadly, the primary goal of sensitivity analysis is to determine which input parameters most influence computational model responses, or deterministic quantities of interest. A ranked list of parameter influences can focus resources for data gathering or model/code development, or can make calibration, optimization, or uncertainty quantification more tractable over a reduced set of parameters. In a post-optimization role, sensitivity information is useful in determining whether or not the response functions are robust with respect to small changes in the optimum design point. The Dakota sensitivity analysis studies recommended in this chapter have important secondary benefits as well: (1) they can help identify key model characteristics such as smoothness, nonlinear trends, and robustness to enable selection of suitable Dakota methods for follow-on studies; and (2) some yield sampling designs that can be used to construct the surrogate models described in Chapter 4 for subsequent analyses.

In the CASL context, a phenomena identification and ranking table (PIRT) might help identify the superset of parameters to consider in a sensitivity analysis study. Then the relative parameter rankings resulting from a Dakota-driven sensitivity study form the basis of a quantitative PIRT, or QPRT. These results could also help prioritize model development or data gathering, or identify insensitive parameters to omit from calibration or UQ studies.

3.1 Terminology

This section introduces key sensitivity analysis terminology and defines the metrics typically used to assign relative ranks to parameter influences on a response.

3.1.1 Local Versus Global Sensitivity

Dakota primarily focuses on sensitivity analysis in a global sense, i.e., over the whole valid parameter domain. We contrast that here with more traditional local or partial derivative-based sensitivity analysis.

Local Sensitivity: In some instances, the term sensitivity analysis is used in a local sense to denote the computation of response derivatives with respect to parameters at a point. These local derivatives can then be used to make design decisions or rank parameter influences. Dakota supports this type of study through numerical finite-differences or retrieval of analytic gradients computed within the analysis code. The desired gradient data is specified in the responses section

of the Dakota input file and the collection of this data at a single point is accomplished through a parameter study method with no steps.

This approach to sensitivity analysis should be distinguished from the activity of augmenting analysis codes to internally compute derivatives using techniques such as direct or adjoint differentiation, automatic differentiation (e.g., ADIFOR), or complex step modifications. These sensitivity augmentation activities are completely separate from Dakota and are outside the scope of this manual. However, once completed, Dakota can utilize these analytic gradients to perform optimization, uncertainty quantification, and related studies more reliably and efficiently. In CASL, some simulation codes such as TSUNAMI have adjoint capabilities and can return not only function value, but derivative data to Dakota, enhancing analyses.

Global Sensitivity: In other instances, the term sensitivity analysis is used in a more global sense to denote the investigation of variability in the response functions over the whole valid range of the input parameters. Dakota supports this type of study through computation of response data sets at a series of sample design points in the parameter space. The series of points is typically defined using a parameter study or a design and analysis of computer experiments (DACE) design, such as orthogonal arrays or space filling Monte Carlo sampling. These more global approaches to sensitivity analysis can be used to obtain trend data even in situations when gradients are unavailable, unreliable, or not indicative of global trends.

This chapter offers guidance solely on Dakota's global sensitivity analysis procedures. Using them typically consists of:

1. Specifying ranges for each parameter and a sensitivity analysis method in the Dakota input
2. Running Dakota which will:
 - (a) construct a sampling design in the parameter hypercube;
 - (b) run the computational model at these points, collecting returned response data; and
 - (c) calculate and output sensitivity metrics to rank inputs.
3. Post-processing the Dakota-generated parameter/response table with external statistics and visualization tools to further assess trends and which input factors most strongly influence the responses

3.1.2 Sensitivity Metrics

Sensitivity metrics output by Dakota are used to assess the relative influence of or rank parameters. The metrics output vary by Dakota method as discussed in Section 3.2, but may include:

- **Correlation coefficients:** Dakota prints correlation tables with the simple (Pearson), partial, and rank (Spearman) correlations between inputs and outputs. These are all bounded between -1 and 1 and measure the strength of the linear relationship between the variables considered. These can be useful to get a quick sense of how correlated the inputs are to each other, and how correlated various outputs are to inputs, but can be misleading for detecting nonlinear relationships. For example a model with a perfectly quadratic input/output relationship centered at zero would have zero correlation hiding the actual strong nonlinear relationship.

The simple correlations are Pearson’s correlation coefficient, which is defined for two factors w and x (where each of these could represent an input or an output) as:

$$\text{Corr}(w, x) = \frac{\sum_i (w_i - \bar{w})(x_i - \bar{x})}{\sqrt{\sum_i (w_i - \bar{w})^2 \sum_i (x_i - \bar{x})^2}}.$$

Partial correlation coefficients are similar, but measure correlation while adjusting for the effects of other variables. For example, in a problem with two inputs and one output where the two inputs are highly correlated, the correlation of the second input and the output may be very low after accounting for the effect of the first input. The rank correlations in Dakota are obtained using Spearman’s rank correlation. Spearman’s rank is the same as the Pearson correlation coefficient except that it is calculated on the rank data. Rank correlation can be more informative when responses vary over orders of magnitude. The correlation analyses are explained further in the Uncertainty Quantification chapter of the Dakota User’s Manual. [1]

- **Morris metrics** [25] are computed from “elementary effects” based on a sample design of large steps around the parameter space. Here each dimension of a M –dimensional input space is uniformly partitioned into p levels, creating a grid of p^M points $x \in \mathbb{R}^M$ at which evaluations of the model $y(x)$ might take place. An elementary effect corresponding to input i is computed by a forward difference

$$d_i(x) = \frac{y(x + \Delta e_i) - y(x)}{\Delta}, \quad (3.1)$$

where e_i is the i^{th} coordinate vector, and the step Δ is typically taken to be large (this is not intended to be a local derivative approximation), e.g., for an input variable scaled to $[0, 1]$, $\Delta = \frac{p}{2(p-1)}$, so the step used to find elementary effects is slightly larger than half the input range.

The distribution of elementary effects d_i over the input space characterizes the effect of input i on the output of interest. After generating N samples from this distribution, their mean,

$$\mu_i = \frac{1}{N} \sum_{j=1}^N d_i^{(j)}, \quad (3.2)$$

modified mean

$$\mu_i^* = \frac{1}{N} \sum_{j=1}^N |d_i^{(j)}|, \quad (3.3)$$

(using absolute value) and standard deviation

$$\sigma_i = \sqrt{\frac{1}{N-1} \sum_{j=1}^N \left(d_i^{(j)} - \mu_i \right)^2} \quad (3.4)$$

are computed for each input i . The mean and modified mean give an indication of the overall effect of an input on the output. Standard deviation indicates nonlinear effects or interactions, since it is an indicator of elementary effects varying throughout the input space.

- **Sobol indices:** Dakota can calculate sensitivity indices through Variance-based Decomposition (VBD). 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 sensitivity index T_i . The main effect sensitivity index corresponds to the fraction of the total uncertainty in the output, Y , that can be attributed to input x_i alone. The total effect sensitivity index corresponds to the fraction of the total 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:

$$S_i = \frac{Var_{X_i}[E(Y|X_i)]}{Var(Y)} \quad (3.5)$$

and

$$T_i = \frac{E_{X_{-i}}[Var(Y|X_{-i})]}{Var(Y)} = \frac{Var(Y) - Var_{X_{-i}}[E(Y|X_{-i})]}{Var(Y)} \quad (3.6)$$

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.

When using VBD, a rough guide is that variables with main effect indices greater than $100/M\%$ are significant as they can be considered to have greater than average effect on output variability, barring higher-order interactions. More details on the calculations and interpretation of the sensitivity indices can be found in [33].

- **Main effects** show the effects of a single variable, averaging across the effect of other input variables. For a full factorial design with each of M inputs taking on p levels, the main effect of input variable x_i is calculated at each level $k = 1, \dots, p$ it takes on as

$$m_i^k = \frac{1}{p^{M-1}} \sum_x y(x|x_i = x_i^k).$$

To calculate main effects with Dakota, one can either use (1) the orthogonal array method from DDACE with the supplementary command `main.effects`, or (2) a grid parameter study, which has to be post-processed to compute main effects in an external statistics tool.

Supplementary Approaches: Running any of the parameter study, design of experiments, or sampling methods allows the user to save the results in a tabular data file, which then can be read into a spreadsheet or statistical package for further analysis. One example of this is the well-known technique of scatter plots, in which the set of samples is projected down and plotted against one parameter dimension, for each parameter in turn. Scatter plots with a uniformly distributed cloud of points indicate parameters with little influence on the results, whereas scatter plots with a defined shape to the cloud indicate parameters which are more significant. Related techniques include analysis of variance (ANOVA) [27] and main effects analysis, in which parameters having the greatest influence on the output are identified from sampling results. Scatter plots and ANOVA may be accessed through import of Dakota tabular results into external statistical analysis programs such as R (<http://www.r-project.org>) and Minitab (<http://www.minitab.com>).

3.2 Recommended Methods

This section summarizes a few recommended Dakota sensitivity analysis methods at a high level, shows input file examples, resulting output, and post-processing/visualization approaches that can help. The choice of method will depend on the analysis goal and available computational budget. We begin with high-level best practices before delving into examples of specific methods.

We almost always recommend starting with simple **centered parameter studies** that yield univariate effects only. Do this first with small perturbations, then large variations that span the parameter space. These simple studies test the model interface, assess the relative smoothness of the response, and assess model robustness over single parameter variations. These studies readily determine the effect of a single parameter in a practical way, as they are automated versions of typical “perturb $\pm 5\%$, $\pm 10\%$ ” studies that analysts manually conduct.

The type of follow-on sensitivity study to conduct depends on simulation budget and goal. The cost for various methods is shown in Table 3.1, together with the key metrics that they yield. Here M is the number of input parameters studied, p a user-specified number of increments or partitions in each variable (often taken to be $p = 3$), N a total number of samples in a single Latin hypercube sampling (LHS) replicate, and k a number of replicates (often $k = 4$) which may be needed to evaluate the formulas for Sobol indices from Section 3.1.2.

Table 3.1: Key sensitivity analysis methods with the metrics they produce, roughly ordered by increasing computational cost. M : number of parameters, p : increments per variables, N : total samples in a single replicate, k : number of replicates.

method	design points	metrics
centered parameter study (Sec. 3.2.1)	$p \times M + 1$	univariate effects
global LHS sampling (Sec. 3.2.3)	$N = 2 \times M$ to $10 \times M$	Pearson, partial, Spearman correlations
PSUADE/Morris (Sec. 3.2.4)	$k \times (M + 1)$, with p odd	elementary effects
VBD/Sobol (none)	$N \times (M + 2)$	Sobol main/total effects
full factorial/grid (Sec. 3.2.2)	p^M	correlations, main effects

A **global LHS sampling** study is the most common follow-on study, ideally with $N = 10 \times M$, but possibly as few as $N = 2 \times M$, samples. Global LHS sampling has the benefit of reuse of the sample points for follow-on surrogate construction. Dakota directly outputs correlation coefficients, which when large can indicate parameters surely influencing the response; useful for inclusion-based screening. It also yields data for constructing scatter plots in post-processing analysis. Parameter effects can be confounded, so it can be hard to extract univariate effects, but one can get a good idea of the effect of joint variation with modest samples.

If point reuse or scatter plot diagnostics are not a primary consideration, **PSUADE/Morris** designs offer more inference power, with a similar cost, to LHS designs by facilitating quantitative detection of nonlinear or interaction effects in addition to main effects.

Variance-based decomposition (VBD/Sobol) analysis can offer even more information. This uses replicate LHS or a surrogate (possibly stochastic polynomial approximations as described in Section 6.1.2) to perform a variance-based decomposition that apportions output variance to input factors. The resulting Sobol indices for main and total effects (described in Section 3.1.2) can be helpful for up front sensitivity analysis, when an “80/20” principle applies: fewer than 20

percent of the parameters explain at least 80 percent of total output variance. In this case, one can quickly screen with relatively few model runs. However, this approach is challenging in the presence of strong interactions, due to the potentially substantial data requirements for accurate inference of Sobol indices. Strong interactions are suggested by differences in partial versus simple correlations or main versus total effects in Sobol indices.

For modest numbers of parameters and reasonable model run cost, one can conduct **full factorial** parameter studies with Dakota's grid/multidimensional parameter study or fractional factorial orthogonal array (OA) designs with DACE OA. These can assess main effects of each parameter, even when considering them jointly. A two-level Plackett-Burman design can be good for extremely slow codes, where the number of runs ($N = M + 1$) is severely limited. However these methods are not currently available in Dakota.

3.2.1 Centered Parameter Study

The centered parameter study executes multiple coordinate-based parameter studies, one per parameter, centered about the specified initial values. This is useful for investigation of function contours in the vicinity of a specific point and is a very common model exploration technique where each parameter is increased and decreased by a fixed increment. A set of widely-spaced points in a centered or multidimensional parameter study could be used to determine whether the response function variation is likely to be unimodal or multimodal. A set of closely-spaced points in a centered parameter study could also be used to assess the smoothness of the response functions in order to select a suitable finite difference step size for optimization/calibration (see an example of their use in Listing 5.1 in Section 5.2.1). After computing an optimum design, a parameter study could also be used for post-optimality analysis in verifying that the computed solution is actually at a minimum or constraint boundary and in investigating the shape of this minimum or constraint boundary. (In a parameter study, one may optionally enable Dakota's numerical gradient estimation to calculate local derivative values at each point in the parameter space, but the results are only published to the Dakota console text output, not the tabular data file.)

Dakota Input: This method requires two settings: (1) **step_vector**, a list of real values, each of which specifies the size of the increment or perturbation for a single variable; and (2) **steps_per_variable**, a list of integers that specifies the number of increments p_i per variable in each of the positive and negative coordinate directions. Centered parameter studies are typically conducted with $p_i = 5$ positive and negative increments of each parameter. The total number of samples required is $N = 1 + \sum_{i=1}^M 2p_i$. **step_vector** specifies absolute variable steps for continuous and discrete range variables, but for studies conducted over integer or real discrete set variables (see "Design Variables" in the Dakota Reference Manual [2]), specifies perturbations in index offsets to select from the possible set values. For example, with initial values of (1.0, 1.0), a **step_vector** of (0.1, 0.1), and a **steps_per_variable** of (2, 2), the center point is evaluated followed by four function evaluations (two negative deltas and two positive deltas) per variable. This set of points in parameter space is shown in Dakota screen output in Figure 3.1 and graphically in Figure 3.2.

Dakota Input for Cantilever: A sample Dakota input file for a centered parameter study with the cantilever beam application is shown in Listing 3.1. Note the previously discussed method controls for **centered_parameter_study** in lines 8–10, that all the variables are active (line 18), and that the evaluations will be saved to the tabular data file specified on line 4. The centered study is fully characterized by the initial values and steps for the variables (lines 21 and 25, 10 and 11).

```

Parameters for function evaluation 1:
      1.0000000000e+00 d1
      1.0000000000e+00 d2
Parameters for function evaluation 2:
      8.0000000000e-01 d1
      1.0000000000e+00 d2
Parameters for function evaluation 3:
      9.0000000000e-01 d1
      1.0000000000e+00 d2
Parameters for function evaluation 4:
      1.1000000000e+00 d1
      1.0000000000e+00 d2
Parameters for function evaluation 5:
      1.2000000000e+00 d1
      1.0000000000e+00 d2
Parameters for function evaluation 6:
      1.0000000000e+00 d1
      8.0000000000e-01 d2
Parameters for function evaluation 7:
      1.0000000000e+00 d1
      9.0000000000e-01 d2
Parameters for function evaluation 8:
      1.0000000000e+00 d1
      1.1000000000e+00 d2
Parameters for function evaluation 9:
      1.0000000000e+00 d1
      1.2000000000e+00 d2

```

Figure 3.1: Dakota output showing function evaluations for a centered parameter study with two positive and two negative steps per variable.

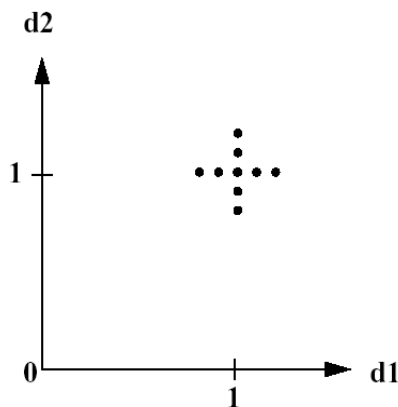


Figure 3.2: Notional example of centered parameter study over two parameters $d1$ and $d2$.

Results and Discussion: The results of the study are depicted in Figure 3.3, where the tabular data generated by Dakota (`cantilever_centered.dat`) has been plotted with Matlab. These plots show that only w and t affect area /weight (the plots for these two variables overlay each other). For the stress and displacement, w and t have the strongest effect, and possibly a nonlinear one as evidence by the curvature in their traces. E and X have a small, but nonzero effect. All input/output relationships appear smooth (no noise or other oscillation is evident). These observations can be verified by studying the equations for the static cantilever problem.

Listing 3.1: Dakota input file showing centered parameter study on the cantilever beam problem.

```

1 strategy,
  single_method
3   tabular_graphics_data
   tabular_graphics_file 'cantilever_centered.dat'
5
method,
7
# do a parameter study in coordinate directions over all 6 parameters
9   centered_parameter_study
   step_vector 0.1 0.1 10 100 10 100
11  steps_per_variable 2
13
variables,
15
# by default, a parameter study won't operate on state parameters
# can change that default behavior be explicitly specifying which
17 # parameters to use (here "all")
   active all
19
   continuous_design = 2
21   initial_point 1.0 1.0
   descriptors 'w' 't'
23
   continuous_state = 4
25   initial_state 40000. 29.E+6 500. 1000.
   descriptors 'R' 'E' 'X' 'Y'
27
interface,
29   direct
   analysis_driver = 'mod_cantilever'
31
responses,
33   num_objective_functions = 3
   response_descriptors = 'area' 'stress' 'displacement'
35   no_gradients
   no_hessians

```

3.2.2 Multidimensional Parameter Study

The multidimensional parameter study computes response data sets for an M -dimensional hyper-grid of input points. This full factorial design is powerful in determining main effects and potential interactions among parameters, but the number of simulation runs quickly becomes prohibitive as

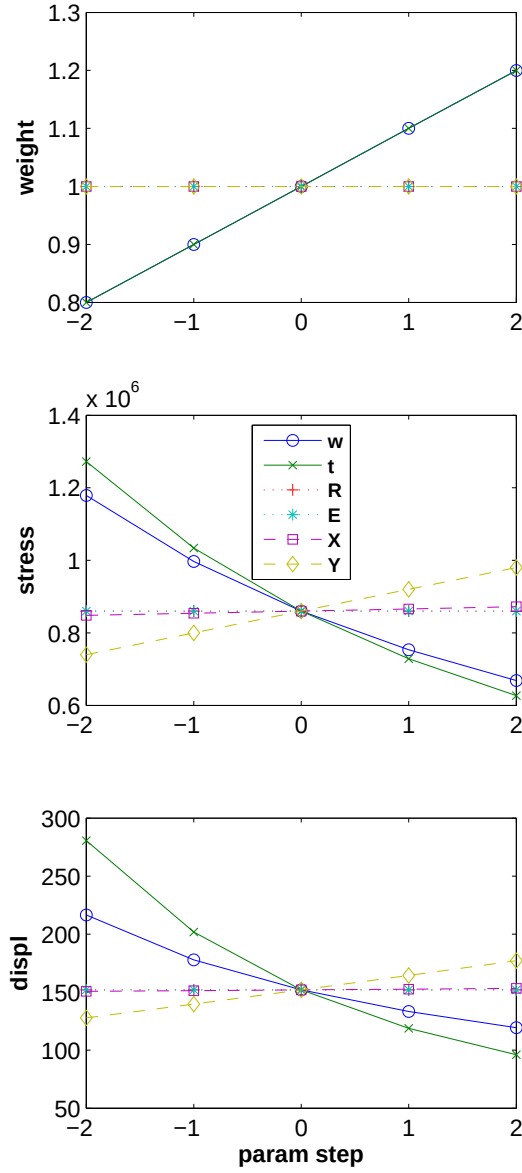


Figure 3.3: Cantilever beam: univariate effects from centered parameter study of each of six parameters on each of three responses.

the dimension of the parameter space M increases. It is presented here mainly because it's easily understandable and tractable/useful for small dimensional problems.

For these studies, each variable is partitioned into equally spaced intervals between its upper and lower bounds, and each combination of the values defined by these partitions is evaluated. The

number of function evaluations performed in the study is:

$$\prod_{i=1}^M (\text{partitions}_i + 1) \quad (3.7)$$

Dakota Input Example: The partitions information is provided using the `partitions` specification, which inputs an integer list of the number of partitions for each variable (i.e., `partitionsi`). Since the initial values will not be used, they need not be specified.

In a two variable example problem with $d1 \in [0,2]$ and $d2 \in [0,3]$ (as defined by the upper and lower bounds from the variables specification) and with `partitions` = (2,3), the interval $[0,2]$ is divided into two equal-sized partitions and the interval $[0,3]$ is divided into three equal-sized partitions. This two-dimensional grid, shown notionally in Figure 3.4, would result in the twelve function evaluations shown in Figure 3.5. See the first example in the Dakota User’s Manual [1]: Tutorial for additional notes to understand this study.

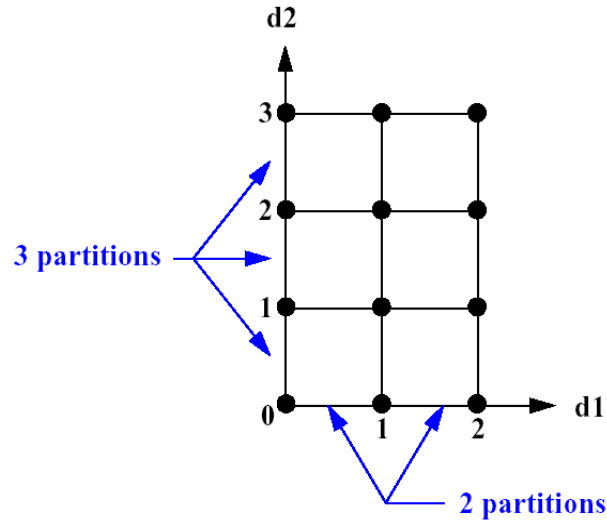


Figure 3.4: Example of multidimensional parameter study.

Dakota Input for Cantilever: Listing 3.2 shows a Dakota input file prescribing a multidimensional parameter study for the cantilever beam problem. On line 10, 9 partitions are specified for w and 6 for t , resulting in $10 \times 7 = 70$ total model evaluations in the (w, t) space. The parameters R , E , X , and Y are held at nominal values using Dakota’s state variable mechanism, as `active all` is commented on line 16. In contrast to the centered parameter study, the active variables are characterized by their lower and upper bounds (lines 19–20).

Results and Discussion: When resources allow a grid parameter study to be conducted, one resulting advantage is that main effects can be calculated. An example is shown in Figure 3.6. In this example, the main effects for w and t are generated by post-processing Dakota’s `cantilever_grid.dat` file using external statistical and plotting software. The left subplot shows the main effect of w , that is the relationship between w and the mean of the displacement, taken over all realization of the other variable t . We observe a smooth, nonlinear effect. Similar is true for the main effect of t in the right subplot.

Listing 3.2: Dakota input file showing grid parameter study on the cantilever beam problem.

```
strategy,
2   single_method
   tabular_graphics_data
4   tabular_graphics_file 'cantilever_grid.dat'

6 method,

8 # conduct grid parameter study with 10 values for width, 7 for thickness
   multidim_parameter_study
10   partitions = 9 6

12 variables,

14 # default is to perform the study over the design variables, leaving
   # state fixed; could override to do all variables with:
16 ##active all

18   continuous_design = 2
       lower_bounds    1.0      1.0
20   upper_bounds      4.0      4.0
       descriptors     'w'      't'
22
       continuous_state = 4
24   initial_state     40000. 29.E+6 500. 1000.
       descriptors     'R'      'E'      'X'      'Y'
26
interface,
28   direct
       analysis_driver = 'mod_cantilever'
30
responses,
32   num_objective_functions = 3
       response_descriptors = 'area' 'stress' 'displacement'
34   no_gradients
       no_hessians
```

```

Parameters for function evaluation 1:
      0.0000000000e+00 d1
      0.0000000000e+00 d2
Parameters for function evaluation 2:
      1.0000000000e+00 d1
      0.0000000000e+00 d2
Parameters for function evaluation 3:
      2.0000000000e+00 d1
      0.0000000000e+00 d2
Parameters for function evaluation 4:
      0.0000000000e+00 d1
      1.0000000000e+00 d2
Parameters for function evaluation 5:
      1.0000000000e+00 d1
      1.0000000000e+00 d2
Parameters for function evaluation 6:
      2.0000000000e+00 d1
      1.0000000000e+00 d2
Parameters for function evaluation 7:
      0.0000000000e+00 d1
      2.0000000000e+00 d2
Parameters for function evaluation 8:
      1.0000000000e+00 d1
      2.0000000000e+00 d2
Parameters for function evaluation 9:
      2.0000000000e+00 d1
      2.0000000000e+00 d2
Parameters for function evaluation 10:
      0.0000000000e+00 d1
      3.0000000000e+00 d2
Parameters for function evaluation 11:
      1.0000000000e+00 d1
      3.0000000000e+00 d2
Parameters for function evaluation 12:
      2.0000000000e+00 d1
      3.0000000000e+00 d2

```

Figure 3.5: Dakota output function evaluations for a grid parameter study, the tensor product of three steps in $d1$ with four steps in $d2$.

The plot shown here was generated using Minitab statistical software, but SAS (<http://www.sas.com>), JMP (www.jmp.com), R (<http://www.r-project.org>), Minitab (<http://www.minitab.com>), Matlab (<http://www.mathworks.com/products/matlab>), and other tools can produce similar graphics. Therefore we recommend CASL analysts use the tools present in the computing environment at their institution. Grid parameter studies also output correlation coefficients, discussed at greater length in the next section.

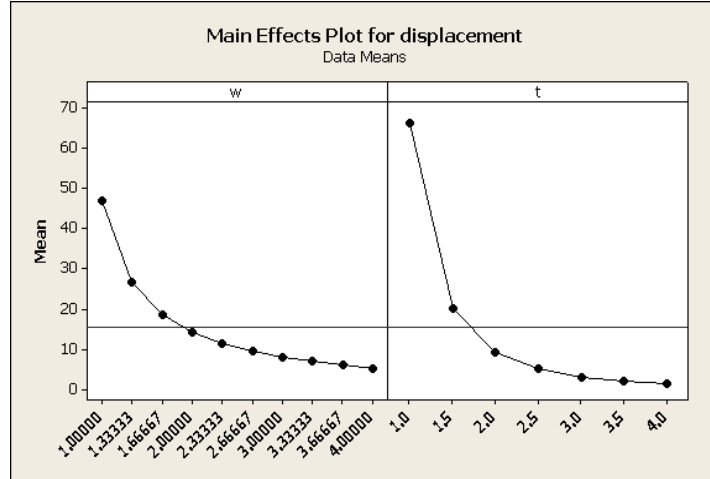


Figure 3.6: Cantilever beam: Minitab-generated main effects for w and t from grid study.

3.2.3 Global LHS Sampling

Monte Carlo sampling methods including Latin hypercube sampling (LHS) are discussed in more detail in Chapter 6, Uncertainty Quantification. For sensitivity analysis with global sampling, variables are typically taken to be uniform on their support. Dakota will generate a space-filling sample design (shotgun blast of points into the M -dimensional parameter space). It will then run the model at these points, and analyze the resulting response data. Here we use a Latin hypercube design for consistency with the UQ recommendations. Latin hypercube designs have better space filling properties and 1-D projections. They will converge statistics at a faster rate than simple Monte Carlo designs.

Dakota Input for Cantilever: A Dakota input file example for conducting a sampling-based study on the cantilever beam is displayed in Listing 3.3. It generates a Latin hypercube design with 100 sample points (lines 11–12) in the six-dimensional parameter space. Note lines 35–38, where the variables are characterized using uniform probability distributions. While not strictly necessary (bounded design variables can work too), this illustrates that Monte Carlo sampling can work with arbitrary probability distributions when needed. The seed on line 29 is specified for study repeatability. If the seed is omitted, Dakota will choose one at random, resulting in a different random design. This can be used to generate replicates to assess variability the statistics. Figure 3.7 shows sets of 2-D projections of one Dakota-generated sampling design.

Results and Discussion: When performing Monte Carlo sampling, Dakota outputs correlations, including the partial correlations shown in the screen output in Figure 3.8. As discussed previously in Section 3.1.2, correlations often give a good quick read on the overall input/output correlation as shown here. Values near 1 indicate strong positive correlation, -1 negative correlation, 0 no correlation. For sensitivity analysis we typically rely on partial correlations as they control for the effect of other variables to provide a more reliable ranking mechanism.

One can visualize the correlation coefficients with external software to more easily see the relative impact. For large numbers of parameters, it is helpful to plot the relative magnitudes (Figure 3.9) or color code with conditional formatting in Excel (Figure 3.10) to more readily differentiate the small from large correlations. Typically, correlations greater than 0.5 in magnitude indicate potentially

Listing 3.3: Dakota input file showing global sampling on the cantilever beam problem.

```

1 strategy,
  single_method
3   tabular_graphics_data
   tabular_graphics_file 'cantilever_sa.dat'
5
method,
7
  # do a sampling-based sensitivity study, optionally with
9  # variance-based decomposition to get higher order sensitivities
  sampling
11   sample_type lhs
   samples = 100
13 #   variance_based_decomp

15 # do a design-of-experiments-based sensitivity using ANOVA to compute
  # the sensitivities
17 # dace oas
  # main_effects
19 # samples = 100

21 # do a Morris One-At-A-Time sensitivity study
  # psuade_moat
23   # must be odd
  # partitions = 3
25   # must be integer multiple of (num_vars + 1)
  # samples = 98
27
  # need these for all methods; seed allows for repeatability
29   seed = 52983

31 variables,

33 # By default DACE and Sampling methods sample over all variables

35   uniform_uncertain = 6
   upper_bounds 48000. 45.E+6 700. 1200. 2.2 2.2
37   lower_bounds 32000. 15.E+6 300. 800. 2.0 2.0
   descriptors 'R' 'E' 'X' 'Y' 'w' 't'
39
interface,
41   direct
   analysis_driver = 'mod_cantilever'
43
responses,
45   num_response_functions = 3
   response_descriptors = 'weight' 'stress' 'displ'
47   no_gradients
   no_hessians

```

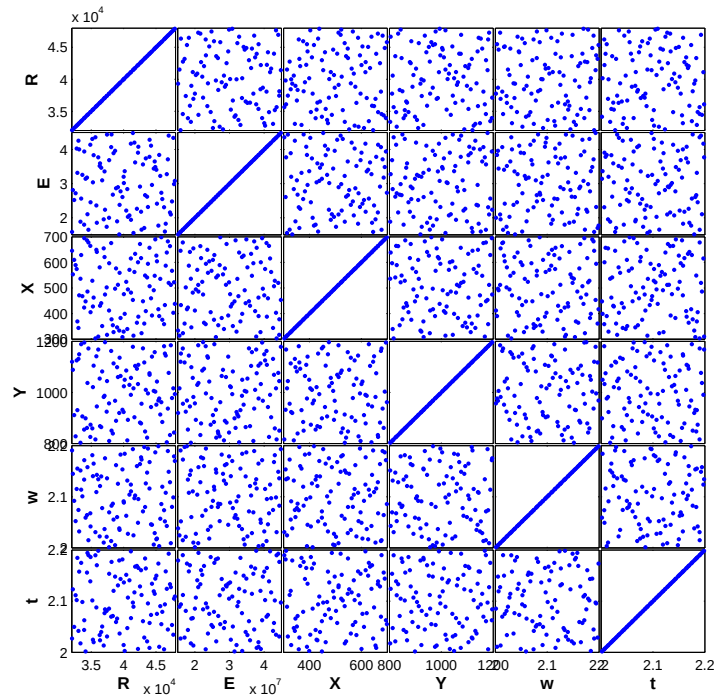


Figure 3.7: Latin hypercube design example: 2-D projections showing locations in the parameter space of 100 Monte Carlo samples using uniform distributions for all six parameters.

```

Partial Correlation Matrix between input and output:
      weight      stress      displ
R  1.36556e-01 -9.89955e-01 -5.82547e-02
E -2.59807e-02  1.51530e-02 -9.53598e-01
X -8.58158e-03  9.96167e-01  3.12725e-01
Y  5.15226e-02  9.96214e-01  7.35493e-01
w  9.99659e-01 -9.84197e-01 -4.20681e-01
t  9.99659e-01 -9.89246e-01 -5.24940e-01

```

Figure 3.8: Dakota output showing partial correlations for the cantilever beam problem.

significant input/output relationships (though specific guidance and interpretation depends on the number of samples, number of variables, and analysis tolerance). Values less than 0.5 should be more carefully studied for possible confounding factors or nonlinearities before discounting their importance.

As correlation coefficients are a linear measure of input/output relationship, it is critical to visualize scatter plots to check for nonlinear trends. Figure 3.11 shows scatter plots generated in Matlab, together with linear regression fits to the data. Some of the scatter plots exhibit nonlinear input/output trends in the cloud of data, for example the plots of E versus displ shows

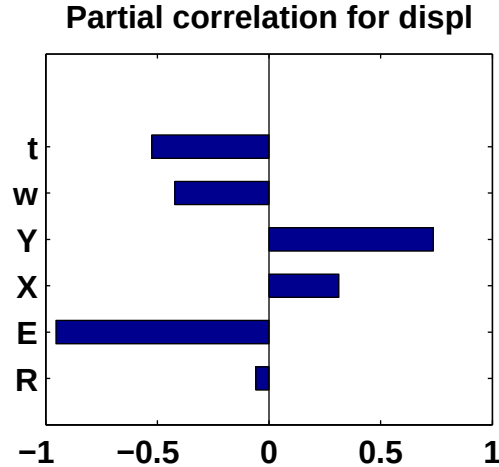


Figure 3.9: Cantilever beam: partial correlation between displacement and inputs.

Partial Correlations for Cantilever			
	weight	stress	displ
R	0.14	-0.99	-0.06
E	-0.03	0.02	-0.95
X	-0.01	1.00	0.31
Y	0.05	1.00	0.74
w	1.00	-0.98	-0.42
t	1.00	-0.99	-0.52

Figure 3.10: Cantilever beam: partial correlation between displacement and inputs, conditionally formatted with Microsoft Excel.

some curvature. Overall, however, the correlations are a good representative indicator of the most important factors for these sample data.

Monte Carlo / LHS sampling for sensitivity analysis screening is typically conducted with number of samples equal to 10 times the number of variables, but budgets often push this down to a factor of two. Also, global Gaussian process models built in Dakota are typically constructed based on a Latin hypercube design. Guidance on the number of samples for constructing such surrogates is provided in Chapter 4.

3.2.4 PSUADE/Morris Method

The Morris One-At-a-Time method, originally proposed by M. D. Morris [25], is a screening method, designed to explore a computational model to distinguish between input variables that have negligible, linear and additive, or nonlinear/interaction effects on the output. The computer experiments performed consist of individually randomized designs which vary one input factor at a time to create a sample of its elementary effects. A more extensive discussion of the method and its metrics can be found in the “Design of Experiments Capabilities” chapter of the Dakota User’s Manual [1].

The file `examples/SensitivityAnalysis/cantilever_morris.in` shows an alternate Dakota

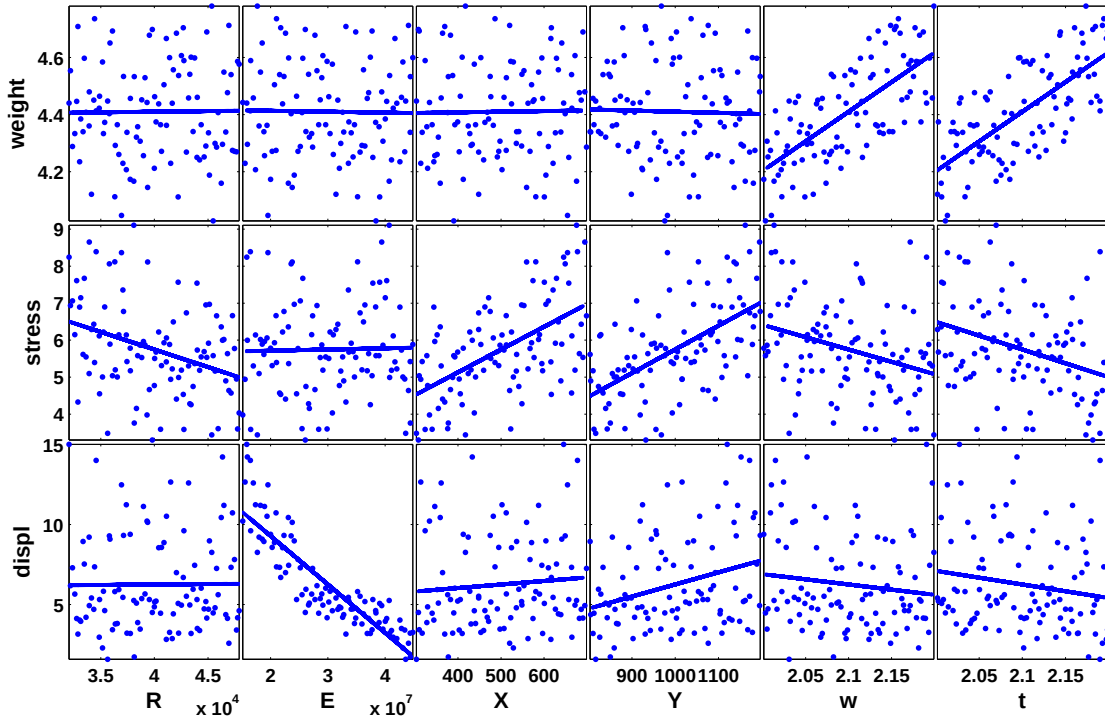


Figure 3.11: Cantilever beam: input/output scatter plots with correlations.

method specification for conducting a Morris screening experiment, resulting in the Morris modified mean and standard deviation of elementary effects, defined above in Section 3.1.2. This method often gives good insight for modest simulation budget. The changed input fragment is:

```
psuade_moat
# must be odd
partitions = 3
# must be integer multiple of (num_vars + 1)
samples = 98
```

In Figure 3.12, the Dakota screen output has been imported into Matlab to plot the modified mean μ^* vs. standard deviation σ of the elementary effects. One can readily observe that for weight, w and t have a strong main effect as expected, and a small nonzero interaction effect. Other variables have no influence. For stress, R has a strong main/linear effect, X and Y have a stronger main effect than w and t , yet w and t have a stronger interaction effect. These are also evident from the cantilever equations. The displacement shows strong main and interaction dependence on E , weaker influence of w, t, X, Y , and no influence of R . This is likely an artifact of the magnitude of E , which is swamping the analysis, and emphasizes the importance of careful input scaling.

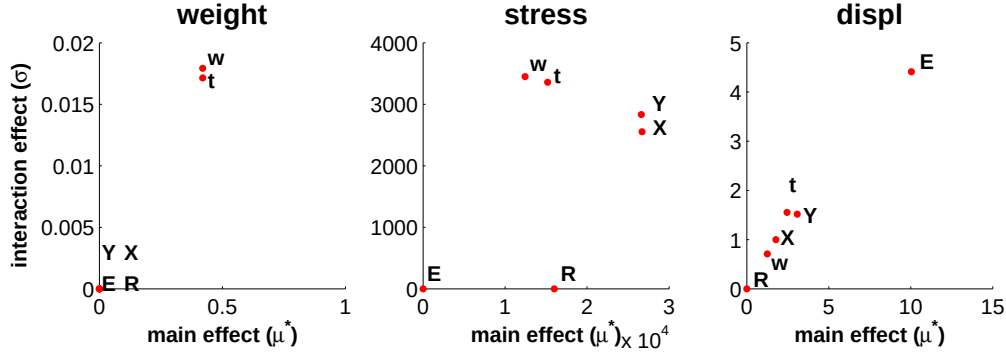


Figure 3.12: Cantilever beam: Morris elementary effects: modified mean (μ^*) and standard deviation (σ).

3.3 Summary and Additional Approaches

Parameter studies, design/analysis of computer experiments (DACE), and general sampling methods share the purpose of exploring the parameter space for sensitivity analysis. When a global space-filling set of samples is desired, then the design of experiments (DOE), DACE, and sampling methods are recommended, with the particular choice depending on computational cost. These techniques are useful for scatter plot and variance analysis as well as surrogate model construction.

We draw a distinction between DOE and DACE methods. DOE are intended for physical experiments containing an element of stochasticity (and therefore tend to place samples at the extreme parameter vertices), whereas the latter are intended for deterministic computer experiments and are more space-filling in nature. Another distinction between DOE/DACE and sampling is drawn based on the distributions of the parameters. DOE/DACE methods typically assume uniform distributions, whereas the sampling (and other uncertainty quantification) approaches in Dakota support a broad range of probability distributions.

Sensitivity analysis method selection recommendations for a broader array of Dakota methods are summarized in Table 3.2. Here are a few highlights to supplement the core recommendations above:

- The vector and list parameter study methods not described here are summarized in the “Parameter Study Capabilities” chapter of the Dakota User’s Manual [1]. List parameter studies run the model at user-specified design points.
- The file `cantilever_sa.in` also demonstrates the use of DACE orthogonal array (lines 17–19) designs with minor changes to the input file. These output other helpful measures of sensitivity described in the “DACE Capabilities” chapter of the Dakota User’s Manual [1].
- When implemented using replicate samples (specifying `method`, `sampling`, `variance_based_decomp` in the Dakota input), variance-based decomposition can require a prohibitive number of model runs. Surrogates such as Gaussian process models (see Section 4.2) are often used to mitigate this cost. Another advanced approach is to use polynomial chaos expansions (Section 6.1.2) which also produce Sobol indices. One may use a structured PCE design generated by Dakota,

or import the points from a previous Monte Carlo sample to build the PCE and quickly calculate the variance-based decomposition.

- Other design types may be more appropriate for polynomial regression such as Box-Behnken [7] or central composite design [8] for quadratic polynomials.

Table 3.2: Guidelines for selection of parameter study, DOE, DACE, and sampling methods.

Method Classification	Applications	Applicable Methods
parameter study	sensitivity analysis, directed parameter space investigations	centered_parameter_study, list_parameter_study, multidim_parameter_study, vector_parameter_study
classical design of experiments	physical experiments (parameters uniformly distributed)	dace (box_behnken, central_composite)
design of computer experiments	variance analysis, space filling designs (parameters uniformly distributed)	dace (grid, random, oas, lhs, oa_lhs), fsu_quasi_mc (halton, hammersley), fsu_cvt, psuade_moat
sampling	space filling designs (parameters have general probability distributions)	sampling (Monte Carlo, LHS) with optional active view override

Chapter 4

Surrogate Models

This chapter introduces the basic theory and use of Dakota’s polynomial regression and kriging surrogate models. Surrogate models are typically employed to provide computationally efficient approximate representations of trends and residual (error) processes in physical data or code output. The terms “emulator,” “response surface,” and “meta-model” refer to the generation of surrogate predictions with associated uncertainty quantification. In the Dakota context, surrogate models are automatically generated based on empirical samples of the true simulation model’s input/output relationship. This type of surrogate can be contrasted with physics-based surrogates which make simplifying assumptions to create a simpler, faster running simulation model.

In the following, $Y(x)$ denotes a surrogate for the physical or computational response of interest $f(x)$, $\tilde{\mu}(x)$ denotes the prediction of $Y(x)$ evaluated at the input x (emulator mean), and $\tilde{\sigma}(x)$ denotes the standard error of prediction evaluated at the input x (emulator standard error). Emulators are constructed from physical data or code output collected on a sample design of N runs in M input dimensions. Good practice requires the run size N to be at least as large as the number of surrogate model parameters requiring estimation (degrees of freedom). Dakota enforces this practice by exiting and returning an error message if N is too small. The degrees of freedom (and thus minimum N) will be stated for each surrogate model introduced in the ensuing sections.

The surrogate models considered in this chapter assume the distribution of residuals (difference between experimental or code output and emulator prediction) is modeled as mean-zero Gaussian. Diagnostics such as a normal probability plot of standardized emulator residuals can be used to check this assumption, as illustrated in Section 7.3.2. Under the assumption of Gaussian residuals, outputs used to construct emulators will then follow a multivariate Gaussian distribution, assuming all unknown surrogate model parameters are fixed. The *likelihood* of observing the given outputs for any given value of the surrogate model parameters is the associated value of this multivariate Gaussian density function. In this chapter, most surrogate model parameters are estimated by *maximum likelihood* in Dakota, by finding a value for these parameters that maximizes the likelihood function. This value is referred to as the maximum likelihood estimate (MLE). Parameter estimation via MLE differs from the Bayesian approach to parameter inference discussed in Uncertainty Quantification, Chapter 6, in that the latter results in a probability distribution for the surrogate model parameters. In practice, the contribution to overall emulator uncertainty induced by surrogate model parameter uncertainty from Bayesian inference is often small compared with the uncertainty arising from the residual error process itself. Computations in Dakota involving surrogate model indirection can thus be sped up considerably by utilizing a point estimate (the

MLE) in place of a probability distribution for the surrogate model parameters.

Section 4.1 considers polynomial surrogates, which are typically used to model observed trends in data from physical experiments resulting from perturbing input parameters that describe physical scenarios of interest. Polynomial surrogates model these trends as a regression relationship that is linear in the unknown coefficients (but not necessarily in the inputs themselves). Experimental errors are assumed to be independently distributed as mean-zero Gaussian distributions having common variance. Polynomial surrogates *smooth* the observed data by finding the best-fitting trend model of those in the class specified by the user.

Section 4.2 considers kriging surrogates, which are typically used to model observed trends in outputs from computationally intensive code runs resulting from input parameter perturbations. In this setting, input parameters describe physical scenarios of interest as well as uncertain initial or boundary conditions, or possibly uncertain closure model parameters that are calibrated to experimental data (see Optimization and Deterministic Calibration, Chapter 5 and Uncertainty Quantification, Chapter 6). Assuming the quantities of interest vary smoothly with input perturbations, the information available from a small number of code runs can be used effectively by the kriging surrogate to infer a correlated error structure. The kriging surrogate is thus able to borrow strength from the small set of code runs conducted to quickly predict code output at user specified input settings with quantified uncertainty. Code runs “closer” to a desired prediction site are weighted more heavily in constructing the kriging emulator than those further away with respect to the inferred correlation structure.

Typically, kriging surrogates *interpolate* the set of code runs made for the purpose of inferring the surrogate model parameters. That is, emulator predictions at input sites corresponding to available code runs are equal to the calculated code outputs with zero uncertainty. This interpolation property is often desirable for emulating deterministic codes that produce a single output value over repeated runs at the same input point. However, as described in Section 4.2 it is possible to relax this interpolation property when fitting a kriging model. This option is useful for leveraging the flexibility of kriging surrogates to model data from physical experiments (e.g. if polynomial surrogates do not provide an adequate fit), or to model output from stochastic codes or deterministic codes subject to numerical or high frequency noise.

As described in the ensuing sections, diagnostics are available to assess the predictive capability of a chosen surrogate model. We synthesize the discussion above to make the following recommendations on the initial choice of surrogate model. Polynomial surrogates are often employed to model outputs from physical experiments. Polynomial or kriging surrogates may be utilized to model code output. Polynomial surrogates are appropriate when very few input parameters are being varied or prior knowledge indicates the trend contains minimal contribution from input nonlinearity and interaction. Subject to these considerations regarding trend complexity, polynomial surrogates are particularly appropriate for modeling stochastic or noisy code outputs. Kriging surrogates are ideal for settings in which weak or nonexistent prior information exists about trend complexity. They are often able to successfully infer the trend from a modest set of training runs, yielding an advantage over the effort often required to successfully fit a suitable polynomial surrogate in the presence of little or no prior information about trends. Kriging surrogates are also capable of fitting nonlinear trends not easily represented by low order polynomial surrogates. For deterministic codes, the uncertainty quantification provided by kriging surrogates conforms to the notion that prediction uncertainty should diminish near input points at which training runs were conducted. This behavior does not occur with polynomial surrogates.

4.1 Polynomial Regression Models

Linear, quadratic, and cubic polynomial surrogate models are available in Dakota. The form of the linear polynomial model is

$$Y(x) = \beta_0 + \sum_{i=1}^M \beta_i x_i + \epsilon(x);$$

the form of the quadratic polynomial model is

$$Y(x) = \beta_0 + \sum_{i=1}^M \beta_i x_i + \sum_{i=1}^M \sum_{j \geq i}^M \beta_{ij} x_i x_j + \epsilon(x);$$

and the form of the cubic polynomial model is

$$Y(x) = \beta_0 + \sum_{i=1}^M \beta_i x_i + \sum_{i=1}^M \sum_{j \geq i}^M \beta_{ij} x_i x_j + \sum_{i=1}^M \sum_{j \geq i}^M \sum_{k \geq j}^M \beta_{ijk} x_i x_j x_k + \epsilon(x).$$

In all of the polynomial models, $Y(x)$ is the response of the polynomial model plus a mean-zero Gaussian error model ($\epsilon(x)$); the x_i, x_j, x_k terms are the components of the M -dimensional input parameter values; and the $\beta_0, \beta_i, \beta_{ij}$, and β_{ijk} terms are the polynomial coefficients. The number of coefficients, N_β , depends on the order of the polynomial model. For the linear polynomial, $N_\beta = M + 1$; for the quadratic polynomial, $N_\beta = (M + 1)(M + 2)/2$; and for the cubic polynomial, $N_\beta = (M^3 + 6M^2 + 11M + 6)/6$. The errors $\epsilon(\cdot)$ associated with each output are assumed to be independently distributed and have constant variance σ^2 .

There must be at least N_β data samples (i.e. $N \geq N_\beta$) in order to form a fully determined linear system and solve for the polynomial coefficients β . In most applications, the linear system will be overdetermined (i.e. $N > N_\beta$). To solve such systems, Dakota employs a least-squares approach involving a QR factorization-based numerical method. Note that $N \geq N_\beta + 1$ is required to estimate both β and the error variance σ^2 .

The maximum likelihood value of β is computed via ordinary least squares,

$$\hat{\beta} = (G^T G)^{-1} G^T y.$$

Here y is the N -vector of observed experimental data and G is a N by N_β matrix that contains evaluations of the polynomial basis functions at all runs in the N by M sample design X , $G_{ij} = g_j(X_i^T)$, $i = 1, \dots, N$; $j = 1, \dots, N_\beta$.

The emulator mean $\tilde{\mu}(x)$ evaluated at input x is a best linear unbiased predictor of the surrogate $Y(x)$,

$$\tilde{\mu}(x) = g^T(x) \hat{\beta}.$$

The emulator variance $\tilde{\sigma}^2(x)$ (here, the mean squared error of prediction) evaluated at input x provides a spatially varying measure of prediction uncertainty,

$$\tilde{\sigma}^2(x) = \widehat{\sigma^2} \left(1 + g^T(x) (G^T G)^{-1} g(x) \right),$$

where an unbiased estimate of σ^2 is

$$\widehat{\sigma^2} = \frac{(y - G\hat{\beta})^T (y - G\hat{\beta})}{N - N_\beta}.$$

The utility of the polynomial models stems from two sources: (1) over a small portion of the parameter space, a low-order polynomial model is often an accurate approximation to the true data trends, and (2) the least-squares procedure provides a surface fit that smooths out noise in the data. For this reason, surrogate-based optimization is often successful when using polynomial models, particularly quadratic models. However, a polynomial surface fit may not be the best choice for modeling data trends over the entire parameter space, unless it is known a priori that the true data trends are close to linear, quadratic, or cubic. Furthermore, in general polynomial models will not *interpolate*, i.e. predict with zero uncertainty, the data they are built from. If interpolation is desired (as with deterministic computational models subject to negligible numerical noise), or if lack of fit is observed, users should consider the kriging emulators described in the following section. See [27] for additional information on polynomial models.

4.1.1 Fitting Polynomial Surrogates in Dakota

In the following discussion, `typewriter` font indicates the names of Dakota input or output files, keywords, commands, or results. All input, output, and log files mentioned in this subsection can be found in `examples/SurrogateModels`. Listings 4.1 and 4.2 together show the Dakota input file `cantilever_polynomial_eval.in` for evaluating a quadratic polynomial emulator of the cantilever beam model outputs at user-specified values of the inputs x . The values of x at which to evaluate the emulator are given in the `list_parameter_study` block of the input file (line 22). The keyword `sampling` (line 58), together with `sample_type lhs` (line 60) tells Dakota to generate a Latin hypercube sample. Dakota then runs the cantilever beam model on the resulting 60 input settings (line 61), producing output for the three indicated response variables (area, stress, and displacement). The design and output results are then used by Dakota to fit a quadratic polynomial model to each response (line 37). Dakota also allows `linear` and `cubic` polynomial trend options. Finally, predicted responses corresponding to each user specified input setting are written to the file `cantilever_polynomial_evals.dat` (line 9):

%eval_id	R	E	X	Y	w	t	area	stress	displacement
1	32887.97269	36889581.78	408.9821034	1148.6638	2.099541047	2.004037063	4.207558073	76546.06547	5.914253488
2	43482.92921	16550393.66	639.6736606	1051.841799	2.014840025	2.19097364	4.414461384	64938.20937	11.96142127
3	36220.72048	44845837.48	561.1743862	935.4138043	2.109386061	2.136922015	4.507593512	57481.41311	2.627795226
4	47325.31136	25228026.72	348.841119	876.590786	2.190414084	2.082472944	4.561478066	28938.90198	5.259963951

By specifying a random number seed (line 59), the results obtained from running Dakota are repeatable, which is useful for regression testing purposes and for situations in which intermediate results (such as experimental designs) must be regenerated exactly in follow-on Dakota analyses.

In this example, we chose to generate a Latin hypercube sample to fit the quadratic surrogate model. Latin hypercube samples are frequently used to fit kriging surrogates (Section 4.2) for complex computational model quantities of interest due to their space-filling nature; however, they can be less efficient than various alternatives more commonly selected for fitting polynomial surrogates, particularly when modeling responses from physical experiments subject to observation error. Two such alternatives for obtaining quadratic emulators are available in Dakota, selected by replacing the `sampling` block with `dace central_composite` for central composite designs [8] or `dace box_behnken` for Box-Behnken designs [7]. Central composite designs generate $N = 1 + 2M + 2^M$ samples, while Box-Behnken designs generate $N = 1 + 4M(M - 1)/2$ samples. A seed can also be specified with `dace` options for repeatability. The option `dace oas` allows for generation

Listing 4.1: Dakota input file producing predictions at user-specified inputs from a quadratic polynomial emulator for the cantilever beam problem.

```

# Build and evaluate a quadratic polynomial emulator of cantilever beam
2 # at a user specified set of points

4 # Top-level controls
strategy
6   single_method
    method_pointer = 'EvalSurrogate'
8   tabular_graphics_data
    tabular_graphics_file = 'cantilever_polynomial_evals.dat'
10
# Method to perform evaluations of the emulator
12 method
    id_method = 'EvalSurrogate'
14    model_pointer = 'SurrogateModel'

16 # Verbose will show the type form of the surrogate model
    output verbose
18
# -----
20 # Emulator evaluation option #1: Provide user specified inputs in the
#                               Dakota input file
22 list_parameter_study
    list_of_points
24     #           R           E           X           Y           w           t
        32887.97269 36889581.78 408.9821034 1148.6638 2.099541047 2.004037063
26     43482.92921 16550393.66 639.6736606 1051.841799 2.014840025 2.19097364
        36220.72048 44845837.48 561.1743862 935.4138043 2.109386061 2.136922015
28     47325.31136 25228026.72 348.841119 876.590786 2.190414084 2.082472944
# -----
30
# Surrogate model specification
32 model
    id_model = 'SurrogateModel'
34    surrogate global
        dace_method_pointer = 'DesignMethod'
36        # Quadratic polynomial model
        polynomial quadratic
38        # compute and print diagnostics after build
        metrics 'rsquared' 'root_mean_squared'
40        press

```

Listing 4.2: (Continued) Dakota input file producing predictions at user-specified inputs from a quadratic polynomial emulator for the cantilever beam problem.

```

42 variables,
    uniform_uncertain = 6
44     upper_bounds      48000.  45.E+6  700.  1200.  2.2  2.2
    lower_bounds      32000.  15.E+6  300.  800.  2.0  2.0
46     descriptors      'R' 'E' 'X' 'Y' 'W' 't'

48 responses
    response_functions = 3
50     descriptors = 'area' 'stress' 'displacement'
    no_gradients
52     no_hessians

54 # Method to generate a 60 run Latin hypercube design to build the emulator
method
56     id_method = 'DesignMethod'
    model_pointer = 'SimulationModel'
58     sampling
        seed = 20
60         sample_type lhs
        samples = 60
62

# The true simulation model to evaluate to build the emulator
64 model
    id_model = 'SimulationModel'
66     single
        interface_pointer = 'SimulationInterface'
68

interface,
70     id_interface = 'SimulationInterface'
    direct
72     analysis_driver = 'mod_cantilever'

```

of orthogonal array designs [18] to support polynomial surrogate fits and main effects analysis (Chapter 3).

Dakota was run with the command `dakota -i cantilever_polynomial_eval.in >& reg.log &`. The log file `reg.log` provides the surrogate parameter estimates $\hat{\beta}$ and $\widehat{\sigma^2}$, as well as leave-one-out cross-validation root mean squared prediction errors (input file, line 40) for each output, which evaluated to `8.3476172226e-16`, `5.2744083954e+01`, and `3.7407615270e-01` for area, stress, and displacement respectively. The root mean squared prediction error (RMSPE) provides a good overall estimate of emulator out-of-sample predictive capability.

A K -fold cross-validation procedure is also available in Dakota by replacing the `press` option (line 40) with `cross_validate folds = K`. The K -fold cross-validation procedure involves randomly partitioning the training data into K subsets, each containing approximately $100/K\%$ of the data. For each subset, the $100/K\%$ of the design runs and their corresponding outputs are “held out”, and the surrogate refitted with the remaining $100 * (K - 1)/K\%$ of the design runs and associated outputs. The $100/K\%$ hold-out points are then predicted by the refitted emulator. Thus a total of K surrogate rebuilds are conducted, each leaving out $100/K\%$ of the original training data, and the RMSPE is computed for all of the hold-out design points over the whole procedure. Running Dakota with $K = 10$, the RMSPE evaluated to `6.1748098826e-16`, `5.3512924399e+01`, and `3.8966016172e-01` for area, stress, and displacement respectively. These results are qualitatively similar to the leave-one-out cross-validation results. Note that the leave-one-out cross-validation procedure is equivalent to N -fold cross-validation.

4.2 Kriging and Gaussian Process Models

The set of techniques known as kriging 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 [9]. 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.

Suppose outputs (e.g. deposits from the borehole sites) are to be collected at the N inputs $\{x_1, \dots, x_N\}$ to train a kriging surrogate $Y(x)$ for the output at input x . The emulator used for prediction of $Y(x)$ is assumed to be linear in these outputs, $\hat{Y}(x) = \sum_{i=1}^N c_i(x) Y(x_i)$. The coefficient vector $c(x)$ is determined by minimizing the mean squared error of prediction $E \left[\left(Y(x) - \hat{Y}(x) \right)^2 \right]$ for fixed correlation parameters, subject to the unbiasedness constraint $E \left[\hat{Y}(x) \right] = E \left[Y(x) \right]$. The notation $E[\cdot]$ indicates expected value. The optimized mean squared error of prediction quantifies the uncertainty in prediction of $Y(x)$.

Kriging surrogates can also be derived from Gaussian processes, which further assume that arbitrary collections of observed outputs follow multivariate Gaussian distributions for fixed trend, variance and correlation parameters. In Dakota, the Gaussian process (GP) framework is adopted to facilitate covariance parameter estimation through techniques such as maximum likelihood and uncertainty quantification via standard statistical inference methods.

A kriging surrogate, $Y(x)$, consists of a trend function (frequently a linear model $g^T(x)\beta$) plus a Gaussian process error model ($\epsilon(x)$) that modifies the trend function locally,

$$Y(x) = g^T(x)\beta + \epsilon(x).$$

This specifies a stochastic process representation of the unknown true surface $f(x)$. The error process $\epsilon(x)$ is assumed initially to have mean zero and constant variance σ^2 . Furthermore, correlated errors in the input space are allowed by specifying a covariance function. In particular, the covariance between the errors at two arbitrary input locations x and x' is modeled as

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

where $r(x, x')$ is a correlation function (i.e. a symmetric, positive definite function satisfying $r(x, x) = 1$). In the following, this correlation function is assumed to depend on the values of unknown parameters ϕ , designated $r(x, x'|\phi)$.

When the true surface results from evaluation of a deterministic computational model, the error process $\epsilon(x)$ can be specified in such a way that the emulator will interpolate, with zero uncertainty, the model runs it was built from. This is accomplished through selection of a correlation function with requisite smoothness properties, discussed further below. The error process specification is modified as follows to accommodate physical experiments subject to measurement or replicate variability, stochastic computational models, or deterministic computational models for which numerical or high frequency noise is of concern,

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

where

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

and Δ^2 is the variance of the observational or numerical error. In the ensuing discussion, the term “nugget” refers to the ratio $\eta = \Delta^2/\sigma^2$.

Figure 4.1 illustrates the behavior of emulators that interpolate versus emulators that smooth output data. In each panel, the red filled circles represent observed outputs and the red curve designates the emulator mean. The gray shaded area outlines the ensemble of pointwise 95% prediction intervals as a function of the input x , computed using the emulator standard error. That is, for any specified x , there is 95% probability that a future sampled output would lie between the lower and upper bounds of the gray shaded area. The left panel illustrates an error process selected for interpolation. Note the emulator mean passes directly through the three observed outputs, and there is no emulator uncertainty at these points (zero width to the prediction intervals). Prediction uncertainty grows as the input level x moves away from any input location at which output is observed. The right panel illustrates an error process selected for smoothing. In this case, the emulator mean defines a smooth surface that predicts the four observed outputs without being required to recover their values exactly with zero uncertainty. The prediction intervals do not narrow as x approaches an input location at which output is observed. Kriging models with no nugget effect ($\eta = 0$) will interpolate, while polynomial models or kriging models with a positive nugget effect ($\eta > 0$) will smooth.

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

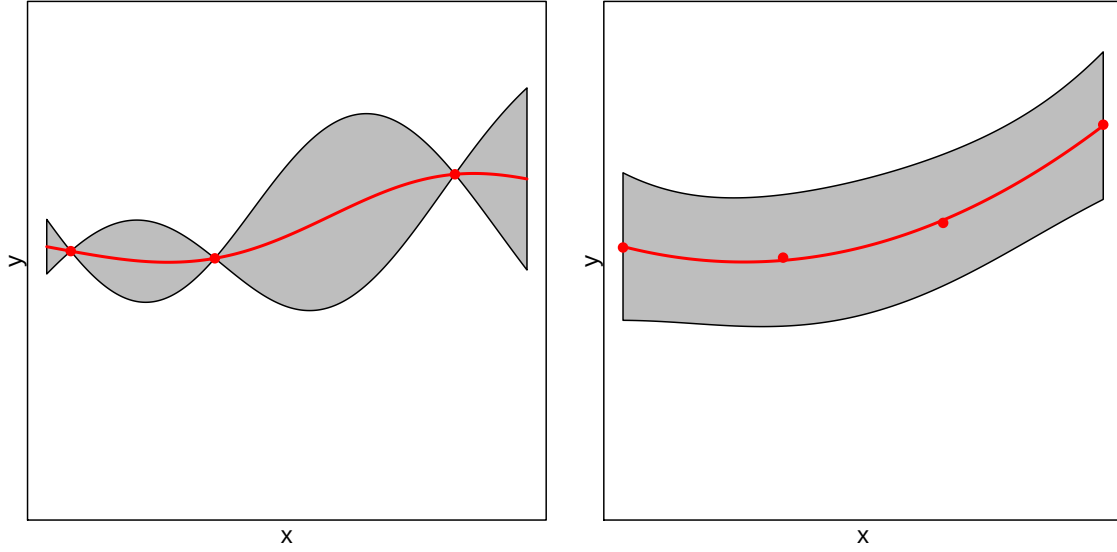


Figure 4.1: Emulator that interpolates (left) and smooths (right) observed output.

treated as a known constant, usually zero, $g^T(x)\beta \equiv 0$. Universal kriging [24] uses a general polynomial trend model $g^T(x)\beta$ with coefficients determined by generalized least squares regression. Dakota allows specification of linear or quadratic trend models. For quadratic models, the user can choose to include or omit the interaction terms among the input variables. Ordinary kriging is essentially universal kriging with a trend order of zero, i.e. the trend function is treated as an unknown constant and $g(x) = 1$. As before, N_β denotes the number of basis functions in $g(x)$ and therefore the number of elements in the vector β . Ordinary kriging is typically selected when emulating code output, as GP models often have sufficient flexibility to detect complex trends without the need for estimating the additional parameters required by a more complex trend function. However, if prior knowledge about more complex trends is available, it should be incorporated into trend function specification, particularly if the emulator is to be used for extrapolation.

The maximum likelihood value of β for fixed correlation parameters is computed via generalized least squares utilizing the correlation matrix $R(\phi)$ of the error process,

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

Here G is a N by N_β matrix that contains evaluations of the polynomial basis functions at all runs in the N by M sample design X , $G_{ij} = g_j(X_i^T)$, $i = 1, \dots, N$; $j = 1, \dots, N_\beta$. The real, symmetric, positive-definite correlation matrix $R(\phi)$ from the error model contains evaluations of the correlation function $r(\cdot, \cdot | \phi)$ at all pairwise combinations of runs in the sample design X ,

$$R_{ij}(\phi) = R_{ji}(\phi) = r(X_i^T, X_j^T | \phi) = r(X_j^T, X_i^T | \phi)$$

There is a single family of correlation functions implemented in Dakota, the power exponential:

$$r(x_i, x_j | \phi) = \exp \left(- \sum_{k=1}^M \phi_k |x_{ik} - x_{jk}|^\gamma \right)$$

where $0 < \gamma \leq 2$ and $\phi_k > 0$. The sample paths of a process equipped with this correlation function are everywhere continuous and nowhere differentiable for $0 < \gamma < 2$, and analytic for $\gamma = 2$. Dakota allows only this latter specification, referred to as the squared exponential or Gaussian correlation function. Note that $N \geq N_\beta + M + 1$ is required to estimate β , the error variance σ^2 , and the correlation function parameters ϕ . Estimation of the nugget η in addition to all other surrogate parameters requires $N \geq N_\beta + M + 2$.

The emulator mean $\tilde{\mu}(x|\phi)$ evaluated at input x is a best linear unbiased predictor of the surrogate $Y(x)$,

$$\tilde{\mu}(x|\phi) = g^T(x)\hat{\beta}(\phi) + \bar{r}^T(x|\phi)R^{-1}(\phi) \left(y - G\hat{\beta}(\phi) \right) ,$$

where $\bar{r}_i(x|\phi) = r(x, X_i^T|\phi)$. This emulator will interpolate the data that the kriging model was built from as long as no nugget is specified or fit and its correlation matrix $R(\phi)$ is numerically non-singular.

The emulator variance $\tilde{\sigma}^2(x|\phi)$ (here, the mean squared error of prediction) evaluated at input x provides a spatially varying measure of prediction uncertainty,

$$\begin{aligned} \tilde{\sigma}^2(x|\phi) = & \widehat{\sigma^2}(\phi) \left(1 - \bar{r}^T(x|\phi)R^{-1}(\phi) \bar{r}(x|\phi) \right) \\ & + \left(g(x) - G^T R^{-1}(\phi) \bar{r}(x|\phi) \right)^T \left(G^T R^{-1}(\phi) G \right)^{-1} \left(g(x) - G^T R^{-1}(\phi) \bar{r}(x|\phi) \right) , \end{aligned}$$

where an unbiased estimate of σ^2 for fixed correlation parameters is

$$\widehat{\sigma^2}(\phi) = \frac{\left(y - G\hat{\beta}(\phi) \right)^T R^{-1}(\phi) \left(y - G\hat{\beta}(\phi) \right)}{N - N_\beta} .$$

Dakota completes construction of the kriging model by using optimization to find a set of correlation parameters ϕ (and if applicable, nugget η) that maximize the likelihood of observing the available data. This is equivalent to minimizing the following objective function,

$$\text{obj}(\phi) = \log \left(\widehat{\sigma^2}(\phi) \right) + \frac{\log(\det(R(\phi))) + \log(\det(G^T R^{-1}(\phi) G))}{N - N_\beta} .$$

Polynomial surrogates are a special case of kriging surrogates, resulting from setting $r(x, x') = 1$ if $x = x'$ and 0 otherwise, and $\Delta = 0$.

4.2.1 Fitting Kriging Surrogates in Dakota

In the following discussion, **typewriter font** indicates the names of Dakota input or output files, keywords, commands, or results. All input, output, and log files mentioned in this subsection can be found in `examples/SurrogateModels`. Listings 4.3 and 4.4 together show the Dakota input file `cantilever_gp_eval.in` for evaluating a GP emulator of the cantilever beam model output at user specified input settings. This surrogate model assumes unknown constant trend (line 35), representing the default ordinary kriging specification. Universal kriging options are also available, by specifying a `linear`, `reduced_quadratic`, or `quadratic` trend. The reduced quadratic option fits a quadratic trend in the absence of interaction terms. This job shows how these input settings can be read from the user specified input file `cantilever_user_points.dat` (line 24), which has the following format:

%eval_id	R	E	X	Y	w	t
1	32887.97269	36889581.78	408.9821034	1148.6638	2.099541047	2.004037063
2	43482.92921	16550393.66	639.6736606	1051.841799	2.014840025	2.19097364
3	36220.72048	44845837.48	561.1743862	935.4138043	2.109386061	2.136922015
4	47325.31136	25228026.72	348.841119	876.590786	2.190414084	2.082472944

The keyword `sampling` (line 56) together with `sample_type lhs` (line 58) tells Dakota to generate a Latin hypercube sample. Dakota then runs the cantilever beam model on the resulting 60 input settings (line 59), producing output for the three indicated response variables (area, stress, and displacement). The design and output results are then used by Dakota to fit a GP with estimated constant trend to each response, and finally predicted responses corresponding to each user specified input setting are written to the file `cantilever_gp_evals.dat` (line 9):

%eval_id	R	E	X	Y	w	t	area	stress	displacement
1	32887.97269	36889581.78	408.9821034	1148.6638	2.099541047	2.004037063	4.209402759	76378.17522	5.360993694
2	43482.92921	16550393.66	639.6736606	1051.841799	2.014840025	2.19097364	4.413155956	65000.4832	12.41942451
3	36220.72048	44845837.48	561.1743862	935.4138043	2.109386061	2.136922015	4.506773804	57529.4144	2.511857535
4	47325.31136	25228026.72	348.841119	876.590786	2.190414084	2.082472944	4.560423301	29135.16863	5.260298986

By specifying a random number seed (line 57), the results obtained from running Dakota are repeatable.

A minimum run size was previously indicated based on the desire to obtain at least as many runs as unknown surrogate parameters. However, more conservative run sizes are often selected, such as a factor of 10 times the number of “active inputs.” Since this latter quantity is often not known, $N = 10M$ is often selected [23]. With $M = 6$, this is the basis for the run size of 60 chosen for this example.

Dakota was run with the command `dakota -i cantilever_gp_eval.in >& gp.log &`. The log file `gp.log` provides the surrogate parameter estimates $\hat{\phi}$, $\hat{\beta}(\hat{\phi})$ and $\widehat{\sigma^2}(\hat{\phi})$, as well as the leave-one-out cross-validation RMSPEs (input file, line 38) for each output, which evaluated to $1.1264992653\text{e-}03$, $1.4913388456\text{e+}02$, and $1.1507679873\text{e-}01$ for area, stress, and displacement respectively. For $K = 10$, the cross-validation RMSPEs for each output evaluated to $1.2088124274\text{e-}03$, $1.4905338700\text{e+}02$, and $1.4178321830\text{e-}01$ for area, stress, and displacement respectively. Kriging fits can be adversely affected if the proportion of runs held out for cross-validation is too large (problem dependent), so it is typical to use $K = N$ (i.e. leave-one-out) in this setting.

Table 4.1 collects the four predictions of each output made by the quadratic polynomial and kriging emulators for comparison with direct code calculations. These predictions track variation in the code calculations well, an observation consistent with the small size of the cross-validation RMSPEs relative to the observed range in the calculations for each output.

4.3 Summary

Table 4.2 summarizes the essential Dakota options for specifying polynomial regression or kriging models, with guidance on which method is to be preferred based on the nature of the physical/computational experiment output and assumptions about the statistical modeling of residual error. Additional details on fitting kriging and other surrogates such as multivariate adaptive regression

Listing 4.3: Dakota input file producing predictions at user specified inputs from a GP emulator with estimated constant trend for the cantilever beam problem.

```

# Build and evaluate a Gaussian process emulator of cantilever beam
2 # at a user specified set of points

4 # Top-level controls
strategy
6   single_method
   method_pointer = 'EvalSurrogate'
8   tabular_graphics_data
   tabular_graphics_file = 'cantilever_gp_evals.dat'
10

12 # Method to perform evaluations of the emulator
method
14   id_method = 'EvalSurrogate'
   model_pointer = 'SurrogateModel'
16

   # Verbose will show the type form of the surrogate model
18   output verbose

20   # -----
   # Emulator evaluation option #2: Provide user specified inputs in a
22   #                               separate file
   list_parameter_study
24   import_points_file 'cantilever_user_points.dat'
   # -----
26

# Surrogate model specification
28 model
   id_model = 'SurrogateModel'
30   surrogate global
   dace_method_pointer = 'DesignMethod'
32   # GP model
   gaussian_process surfpack
34   trend
   constant
36   # compute and print diagnostics after build
   metrics 'rsquared' 'root_mean_squared'
38   press

```

Listing 4.4: (Continued) Dakota input file producing predictions at user specified inputs from a GP emulator with estimated constant trend for the cantilever beam problem.

```

40 variables,
    uniform_uncertain = 6
42     upper_bounds    48000.  45.E+6  700.  1200.  2.2  2.2
    lower_bounds    32000.  15.E+6  300.  800.  2.0  2.0
44     descriptors    'R' 'E' 'X' 'Y' 'W' 't'

46 responses
    response_functions = 3
48     descriptors = 'area' 'stress' 'displacement'
    no_gradients
50     no_hessians

52 # Method to generate a 60 run Latin hypercube design to build the emulator
method
54     id_method = 'DesignMethod'
    model_pointer = 'SimulationModel'
56     sampling
        seed = 20
58         sample_type lhs
        samples = 60
60

    # The true simulation model to evaluate to build the emulator
62 model
    id_model = 'SimulationModel'
64     single
        interface_pointer = 'SimulationInterface'
66

interface,
68     id_interface = 'SimulationInterface'
    direct
70     analysis_driver = 'mod_cantilever'

```

Table 4.1: Polynomial and kriging emulator predictions compared with code calculations.

Area			
Case	Calculation	Polynomial	Kriging
1	4.207558073	4.207558073	4.209402759
2	4.414461384	4.414461384	4.413155956
3	4.507593512	4.507593512	4.506773804
4	4.561478066	4.561478066	4.560423301

Stress			
Case	Calculation	Polynomial	Kriging
1	76625.08025	76546.06547	76378.17522
2	64919.07096	64938.20937	65000.4832
3	57457.96055	57481.41311	57529.4144
4	28991.46319	28938.90198	29135.16863

Displacement			
Case	Calculation	Polynomial	Kriging
1	5.495297702	5.914253488	5.360993694
2	12.52267129	11.96142127	12.41942451
3	2.506561572	2.627795226	2.511857535
4	5.213217663	5.259963951	5.260298986

Table 4.2: Guidelines for selection of surrogate methods.

Method Classification	Applications	Applicable Methods
polynomial regression	smooth fit to physical experiment response or stochastic/noisy computational experiment response [iid residual errors]	polynomial linear, polynomial quadratic, polynomial cubic
kriging	interpolation of deterministic smooth computational experiment response (specify trend option only) or smooth fit to stochastic/noisy computational experiment response (specify trend and nugget options) [correlated residual errors]	gaussian_process surfpack, trend constant, trend linear, trend reduced_quadratic, trend quadratic, nugget ETA (> 0), find_nugget

splines (MARS), simple artificial neural networks, or basic radial basis functions can be found in the Surfpack User's Manual [10].

Chapter 5

Optimization and Deterministic Calibration

The objective of optimization algorithms is to minimize (or maximize) an objective function, typically calculated by the user simulation code, subject to constraints on design variables and responses. Examples of optimization goals include:

- Identify system designs with maximal performance; e.g., case geometry that minimizes drag and weight, yet is sufficiently strong and safe.
- Determine operational settings that maximize system performance, e.g., fuel re-loading pattern yielding the smoothest nuclear reactor power distribution while maximizing output.
- Identify minimum-cost system designs/operational settings, e.g., delivery network that minimizes cost while also minimizing environmental impact.
- Identify best/worst case scenarios, e.g., impact conditions that incur the most damage.
- Calibration: Determine parameter values that maximize agreement between simulation response and target response.

The last goal is a critical use case for CASL. The calibration (parameter estimation, inverse problem) process involves adjusting input parameters to optimally fit a model to experimental or high-fidelity computational model data, find operational settings that best match a prescribed performance profile, or determine source terms for an observed phenomenon. Any Dakota optimization method can be applied to calibration problems, though some are tailored to efficiently address local least squares problem formulations. This chapter emphasizes deterministic model calibration, while non-deterministic approaches such as Bayesian methods are treated in Uncertainty Quantification, Chapter 6.

Available optimization approaches in Dakota include well-tested, proven gradient-based, derivative-free local, and global methods for use in science and engineering design applications. Dakota also offers more advanced algorithms, e.g., to manage multi-objective optimization or perform surrogate-based minimization (useful for noisy or expensive problems). A more extensive treatment of these can be found in “Optimization Capabilities,” “Nonlinear Least Squares Capabilities,” and “Surrogate-Based Minimization” chapters of the Dakota User’s Manual [1]. This chapter continues by introducing optimization terminology needed to select from the available approaches.

5.1 Terminology and Problem Formulations

This section provides a basic introduction to the mathematical formulation of optimization problems. The primary goal of this section is to introduce terms relating to these topics, and is not intended to be a description of theory or numerical algorithms. For further details, consult [4], [14], [16], [28], and [38].

A general optimization problem is formulated as follows:

$$\begin{array}{llll}
 \text{minimize:} & f(x) & \text{objective function} \\
 \text{over:} & x \in \mathbb{R}^M & \text{design variables} \\
 \text{subject to:} & g_L \leq g(x) \leq g_U & \text{nonlinear inequality constraints} \\
 & h(x) = h_t & \text{nonlinear equality constraints} \\
 & a_L \leq A_i x \leq a_U & \text{linear inequality constraints} \\
 & A_e x = a_t & \text{linear equality constraints} \\
 & x_L \leq x \leq x_U & \text{bound constraints}
 \end{array}$$

In this formulation, $x = [x_1, x_2, \dots, x_M]$ is an M -dimensional vector of real-valued *design variables* or *design parameters*. The M -dimensional vectors x_L and x_U , are the lower and upper bounds, respectively, on the design parameters. These bounds define the allowable values for the elements of x , and the set of all allowable values is termed the *design space* or the *parameter space*. A *design point* or a *sample point* is a particular set of values within the parameter space.

The optimization goal is to minimize the *objective function*, $f(x)$, while satisfying the constraints. Constraints can be categorized as either linear or nonlinear and as either inequality or equality. The *nonlinear inequality constraints*, $g(x)$, are “2-sided,” in that they have both lower and upper bounds, g_L and g_U , respectively. The *nonlinear equality constraints*, $h(x)$, have target values specified by h_t . The linear inequality constraints create a linear system $A_i x$, where A_i is the coefficient matrix for the linear system. These constraints are also 2-sided as they have lower and upper bounds, a_L and a_U , respectively. The linear equality constraints create a linear system $A_e x$, where A_e is the coefficient matrix for the linear system and a_t are the target values. The constraints partition the parameter space into feasible and infeasible regions. A design point is said to be *feasible* if and only if it satisfies all of the constraints. Correspondingly, a design point is said to be *infeasible* if it violates one or more of the constraints.

Many different methods exist to solve the optimization problem given by Equation 5.1, all of which iterate on x in some manner. That is, an initial value for each parameter in x is chosen, the *response quantities*, $f(x)$, $g(x)$, $h(x)$, are computed, often by running a simulation, and some algorithm is applied to generate a new x that will either reduce the objective function, reduce the amount of infeasibility, or both. To facilitate a general presentation of these methods, three criteria will be used in the following discussion to differentiate them: optimization problem type, search goal, and search method.

The **optimization problem type** can be characterized both by the types of constraints present in the problem and by the linearity or nonlinearity of the objective and constraint functions. For constraint categorization, a hierarchy of complexity exists for optimization algorithms, ranging from simple bound constraints, through linear constraints, to full nonlinear constraints. By the nature of this increasing complexity, optimization problem categorizations are inclusive of all constraint types up to a particular level of complexity. That is, an *unconstrained problem* has no constraints, a *bound-constrained problem* has only lower and upper bounds on the design parameters, a *linearly-*

constrained problem has linear constraints (and optionally bound constraints), and a *nonlinearly-constrained problem* may contain the full range of nonlinear, linear, and bound constraints (though may omit linear or bound constraints if not applicable). If all of the linear and nonlinear constraints are equality constraints, then this is referred to as an *equality-constrained problem*, and if all of the linear and nonlinear constraints are inequality constraints, then this is referred to as an *inequality-constrained problem*.

Further categorizations can be made based on the linearity of the objective and constraint functions. A problem where the objective function and all constraints are linear is called a *linear programming (LP) problem*. These types of problems commonly arise in scheduling, logistics, and resource allocation applications. Likewise, a problem where at least some of the objective and constraint functions are nonlinear is called a *nonlinear programming (NLP) problem*. These NLP problems predominate in engineering applications and are the primary focus of Dakota.

The **search goal** refers to the ultimate objective of the optimization algorithm, i.e., either global or local optimization. In *global optimization*, the goal is to find the design point that gives the lowest feasible objective function value over the entire parameter space. In contrast, in *local optimization*, the goal is to find a design point that is lowest relative to a “nearby” region of the parameter space. In almost all cases, global optimization will be more computationally expensive than local optimization. Thus, the user must choose an optimization algorithm with an appropriate search scope that best fits the problem goals and the computational budget.

The **search method** refers to the approach taken in the optimization algorithm to locate a new design point that has a lower objective function or is more feasible than the current design point. The search method can be classified as either *gradient-based* or *nongradient-based*. In a gradient-based algorithm, gradients of the response functions are computed to find the direction of improvement. The Hessian (matrix of second derivatives of objectives and constraints with respect to parameters) can also be used in these methods to identify curvature to identify local minima from maxima. Gradient-based optimization is the search method that underlies many efficient local optimization methods. However, a drawback to this approach is that gradients can be computationally expensive, inaccurate, or even nonexistent (the situation for Hessians is typically even worse). In such situations, nongradient-based search methods may be useful. There are numerous approaches to nongradient-based optimization. Some of the more well known of these include pattern search methods (nongradient-based local techniques) and genetic algorithms (nongradient-based global techniques).

5.1.1 Special Considerations for Calibration

Any Dakota optimization algorithm can be applied to calibration problems arising in parameter estimation, system identification, and test/analysis reconciliation. However, nonlinear least-squares methods are optimization algorithms that exploit the special structure of a least squares or sum-of-squares objective function [14]. Here the misfit between vectors of model responses and simulation data is measured in the Euclidean or two-norm.

To exploit the problem structure, more granularity is needed in the response data than is required for a typical optimization problem. That is, rather than using the sum-of-squares objective function and its gradient, least-squares iterators require each term used in the sum-of-squares formulation along with its gradient. This means that the functions in the Dakota response data set consist of the N individual least-squares terms along with any nonlinear inequality and equality constraints. These individual terms are often called *residuals* when they denote differences of

observed quantities from values computed by the model whose parameters are being estimated.

The enhanced granularity needed for nonlinear least-squares algorithms allows for simplified computation of an approximate Hessian matrix. In Gauss-Newton-based methods for example, the true Hessian matrix is approximated by neglecting terms in which residuals multiply Hessians (matrices of second partial derivatives) of residuals, under the assumption that the residuals are zero in expected value under the nonlinear least-squares model. As a result, residual function value and gradient information (first-order information) is sufficient to define the value, gradient, and approximate Hessian of the sum-of-squares objective function (second-order information).

In practice, least-squares solvers will tend to be significantly more efficient than general-purpose optimization algorithms when the Hessian approximation is a good one, i.e., when the neglected component has negligible effect at the solution. Specifically, they can exhibit the quadratic convergence rates of full Newton methods, even though only first-order information is used. Gauss-Newton-based least-squares solvers may experience difficulty when the residuals at the solution are significant. Dakota has three solvers customized to take advantage of the sum of squared residuals structure in this problem formulation. Least squares solvers may experience difficulty when the residuals at the solution are significant, although experience has shown that Dakota's NL2SOL method can handle some problems that are highly nonlinear and have nonzero residuals at the solution.

Specialized least squares solution algorithms can exploit the structure of a sum-of-squares objective function for problems of the form:

$$\begin{array}{llll}
\text{minimize:} & f(\theta) = \sum_{i=1}^N [T_i(\theta)]^2 = \sum_{i=1}^N [y_i(\theta) - d_i]^2 & \text{least squares objective function} \\
\text{over:} & \theta \in \mathfrak{R}^M & \text{calibration variables} \\
\text{subject to:} & g_L \leq g(x) \leq g_U & \text{nonlinear inequality constraints} \\
& h(\theta) = h_t & \text{nonlinear equality constraints} \\
& a_L \leq A_i \theta \leq a_U & \text{linear inequality constraints} \\
& A_e \theta = a_t & \text{linear equality constraints} \\
& \theta_L \leq \theta \leq \theta_U & \text{bound constraints}
\end{array}$$

where $f(\theta)$ is the objective function to be minimized and $T_i(\theta)$ is the i -th least squares term. The bound, linear, and nonlinear constraints are the same as described previously for (5.1). Specialized least squares algorithms are generally based on the Gauss-Newton approximation. When differentiating $f(\theta)$ twice, terms of $T_i(\theta)T_i''(\theta)$ and $[T_i'(\theta)]^2$ result. Because $T_i(\theta)$ is zero in expected value under the nonlinear least-squares model, the Hessian matrix of second derivatives of $f(\theta)$ can be approximated using only first derivatives of $T_i(\theta)$. As a result, Gauss-Newton algorithms exhibit quadratic convergence rates near the solution for those cases when the Hessian approximation is accurate, i.e. the neglected component has negligible effect at the solution. Thus, by exploiting the structure of the problem, the second order convergence characteristics of a full Newton algorithm can be obtained using only first order information from the least squares terms. For problems with nonsmooth gradients or poor finite difference approximations, see [21].

A common example for $T_i(\theta)$ might be the difference between experimental data and model predictions for a response quantity at a particular location and/or time step, i.e.:

$$T_i(\theta) = y_i(\theta) - d_i$$

where $y_i(\theta)$ is the response quantity predicted by the model and d_i is the corresponding experimental data. In this case, θ would have the meaning of model parameters which are not precisely known

and are being calibrated to match available data. This class of problem is known by the terms parameter estimation, system identification, model calibration and test/analysis reconciliation, for example.

This overview of optimization problem formulations and goals approaches underscores that no single optimization method or algorithm works best for all types of optimization problems. The following section offers some basic guidelines for choosing one for specific optimization problems.

5.2 Recommended Methods

In selecting an optimization method, important considerations include the type of variables in the problem (continuous, discrete, mixed), whether a global search is needed or a local search is sufficient, and the required constraint support (unconstrained, bound constrained, or generally constrained). Less obvious, but equally important, considerations include the efficiency of convergence to an optimum (i.e., convergence rate) and the robustness of the method in the presence of challenging design space features (e.g., nonsmoothness). Sensitivity analysis (described in Chapter 3) is a critical precursor to assess problem characteristics prior to choosing and applying an optimization method.

For example the cantilever beam optimization problem posed in (2.3) in Section 2.1 has continuous design variables only, nonlinear inequality constraints, and bound constraints. Sensitivity analysis in Section 3.2 indicated that the objective and constraints are smooth functions of the design variables. This can also be directly inferred from the algebraic physics equations for the cantilever beam (2.1), though cannot for a general physics simulation.

Table 5.1 highlights a few key Dakota optimization approaches and problems for which they are suited. The following sections offer more details on the approaches and input file examples. Without considering specific problem knowledge or characteristics, a derivative-free local pattern search approach is typically a good starting point. If it doesn't find good solutions, move to a genetic algorithm. If it is too costly, move to a surrogate-based approach. However, when problems are smooth and not too multimodal, a local algorithm will outperform these other approaches.

5.2.1 Gradient-Based Local Methods

Gradient-based optimizers are best suited for efficient navigation to a local minimum in the vicinity of the initial point. They are not intended to find global optima in nonconvex design spaces. For global optimization methods, see Section 5.2.3. Gradient-based optimization methods are highly efficient, with the best convergence rates of all of the local optimization methods, and are the methods of choice when the problem is smooth, unimodal, and well-behaved. However, these methods can be among the least robust when a problem exhibits nonsmooth, discontinuous, or multimodal behavior. Figure 5.1 depicts a multimodal function on which gradient-based optimizers will typically find only a nearby local minimum. The derivative-free methods described in Section 5.2.2 are more appropriate for problems with some of these characteristics.

Newton methods are representative of gradient-based optimization methods. These can be derived by applying Newton's method for root finding to $\nabla f(x) = 0$ to find a local minimum of the objective function. The resulting progression from current iterate x_n to next iterate x_{n+1} is then

$$x_{n+1} = x_n - (\nabla^2 f(x_n))^{-1} \nabla f(x_n).$$

Table 5.1: Guidance for selecting from the top recommended Dakota optimization algorithms.

algorithm type / Dakota method	variable type	cost (samples)	goal and characteristics
gradient-based local/ OPT++ Quasi-Newton	continuous	low/ medium	single local solution/improvement, as- sumes smooth input/output mapping
<i>local calibration</i> / OPT++ Gauss Newton	continuous	low/ medium	same as previous line, but tailored to least-squares calibration
derivative-free local / Coliny Pattern Search	continuous	medium/ high	single local solution; better when can't estimate derivatives
<i>local w/surrogate</i> / Surrogate-based Local	continuous	medium	same as “derivative-free local,” but for noisier or more expensive simulations
global / Coliny Evolutionary Algorithm	continuous or discrete	high	global optimality, with ranked family of best solutions
<i>global w/surrogate</i> / Efficient Global	continuous	medium	same as “global,” but for more expen- sive simulations

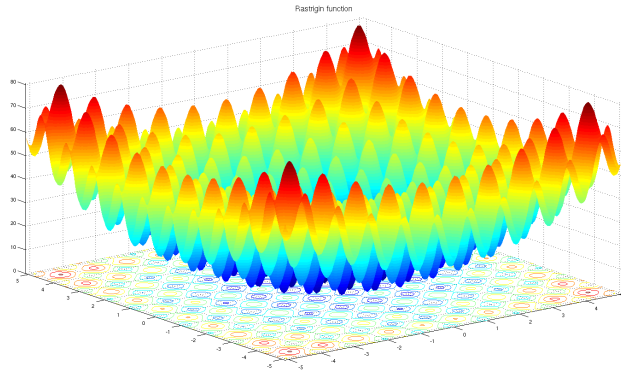


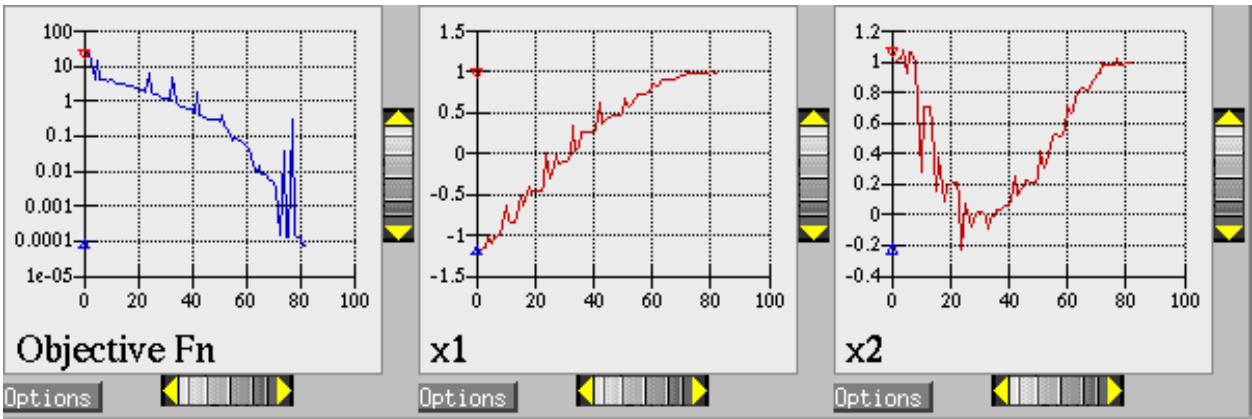
Figure 5.1: Surface plot with contours of an example function that is locally smooth, but globally multimodal.

This naive iteration directly requires the gradient and Hessian ($\nabla^2 f(x)$) of the objective function. In practice this basic iteration is enhanced with strategies to choose an appropriate step length to achieve expected decrease with each iteration, approximate the action of the Hessian-vector product when not directly available, and handle nonlinear constraints, for example via penalty methods.

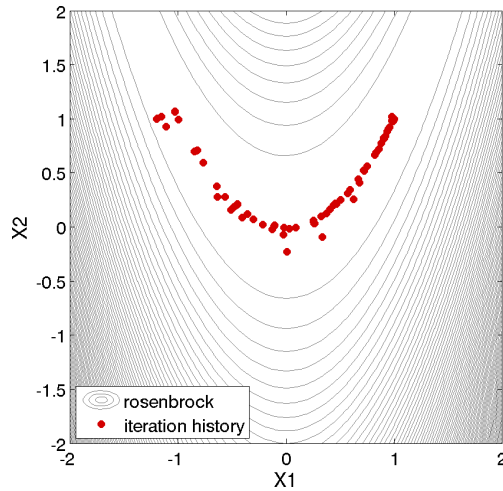
Gradient accuracy is a critical factor for gradient-based optimizers, as inaccurate derivatives will often lead to failures in the search or premature termination of the method. Analytic gradients and Hessians are ideal but often unavailable. If analytic gradient and Hessian information can be provided by an application code, a full Newton method will achieve quadratic convergence rates near the solution. If only gradient information is available and the Hessian information is approximated from an accumulation of gradient data, the superlinear convergence rates can be obtained. It is most often the case for engineering applications, however, that a finite difference

method will be used by the optimization algorithm to estimate gradient values. Dakota allows the user to select the step size for these calculations, as well as choose between forward-difference and central-difference algorithms. The finite difference step size should be selected as small as possible, to allow for local accuracy and convergence, but not so small that the steps are “in the noise.” This requires an assessment of the local smoothness of the response functions using, for example, a parameter study method. Central differencing will generally produce more reliable gradients than forward differencing but at roughly twice the expense.

A typical iteration history and search path for a gradient-based optimization solver is shown in Figure 5.2. Notice the algorithm efficiently going downhill with respect to the contours of the notional objective function.



(a)



(b)

Figure 5.2: Gradient-based unconstrained optimization example: (a) screen capture of the Dakota graphics and (b) sequence of design points (dots) evaluated (line search points omitted).

Recommended method: The recommended local derivative-based solver in Dakota comes from the OPT++ package. A Dakota input example for the cantilever beam optimization problem (see Section 2.1) is shown in Listing 5.1. The iteration starts at the initial iterate for (w, t) specified

in `initial_point` on line 16. The state variables are held fixed at their given values. The algorithm seeks to minimize the cantilever objective function (area), within the bound constraints specified on lines 15 and 17, subject to constraints on stress and displacement specified on line 30. The algorithm will terminate when the convergence criteria from lines 7 and 8 are met. Notice the responses section specifies numerical gradients (line 32), indicating that Dakota should approximate the derivatives of model responses with respect to parameters via finite differences. If the response is a rough or strongly nonlinear function of the parameters, these approximations can be poor and yield bad performance or results. Figure 5.3 shows an excerpt of the Dakota output, showing the optimal design point found.

Listing 5.1: Dakota input file showing local gradient-based optimization on the cantilever beam problem.

```

strategy
2   single_method
    tabular_graphics_data
4     tabular_graphics_file 'cantilever_opt_grad.dat'

6 method
    optpp_q_newton
8     convergence_tolerance 1.0e-4
    constraint_tolerance 1.0e-1

10
11 model
12   single

14 variables
    continuous_design = 2
16     upper_bounds 4.0      4.0
    initial_point 2.4      3.5
18     lower_bounds 1.0      1.0
    descriptors 'w'      't'
20   continuous_state = 4
    initial_state 40000. 29.E+6 500. 1000.
22     descriptors 'R'      'E'      'X'      'Y'

24 interface
    direct
26     analysis_driver = 'mod_cantilever'

28 responses
    objective_functions = 1
30   # constraints assumed <= 0 unless bounds given
    nonlinear_inequality_constraints = 2
32   descriptors = 'area' 'stress' 'displacement'
    numerical_gradients forward
34     fd_step_size 1.0e-6
    no_hessians

```

Variation for calibration: Listing 5.2 shows the Dakota input variation to directly treat calibration with a least-squares specific gradient-based solver. This example tunes the active variables (line 14) $\theta = (E, w, t)$ to match synthetic experimental data from a file, with (line 30) or

```

<<<<< Function evaluation summary: 18 total (18 new, 0 duplicate)
<<<<< Best parameters          =
                2.4000000000e+00 w
                3.5000000000e+00 t
                4.0000000000e+04 R
                2.9000000000e+07 E
                5.0000000000e+02 X
                1.0000000000e+03 Y
<<<<< Best objective function =
                8.4000000000e+00
<<<<< Best constraint values   =
                -4.7108843537e+03
                -2.9685051703e-01

```

Figure 5.3: Dakota output showing optimal local gradient-based optimization result for cantilever beam.

without (line 31) added noise. The data were generated using the following tuning parameter values: $E = 2.85\text{E}7$, $w = 2.5$, $t = 3.0$, and fixed state parameter values: $R = 40000$, $X = 500$, $Y = 1000$. The input file specifies `calibration_terms` instead of `objective_functions` at line 29, indicating to Dakota that it should treat these responses as terms in a least-squares calibration problem. The data here is a set of three observations, one each for area, stress, and displacement. The NL2SOL solver is used for this example as it performs better than the default-recommended OPT++ Gauss-Newton solver.

The calibrated parameter values from the Dakota output are shown in Figure 5.4. Dakota has recovered the true values of the parameters, verifying the operation of the algorithm. The output from local calibration methods also includes confidence intervals on the parameters. With 95% confidence the true value of each parameter lies in the interval specified, given the misfit between the model and corresponding data. When this problem is exercised with noisy data, the confidence intervals expectedly grow larger.

For problems not suitable for local gradient-based optimization, any of the optimization methods discussed in the following sections can be applied to Dakota responses with `calibration_terms`. Dakota will automatically compute the objective function as the sum of squared residuals for presentation to the optimization algorithm.

5.2.2 Derivative-Free Local Methods

Derivative-free methods can be more robust and more inherently parallel than gradient-based approaches. They can be applied in situations where gradient calculations are too expensive or unreliable. In addition, some derivative-free methods can be used for global optimization which gradient-based techniques (see 5.2.1), by themselves, cannot. For these reasons, derivative-free methods are often go-to methods when the problem may be nonsmooth, multimodal, or poorly behaved. It is important to be aware, however, that they exhibit much slower convergence rates for finding an optimum, and as a result, tend to be much more computationally demanding than gradient-based methods. They often require from several hundred to a thousand or more function

Listing 5.2: Dakota input file showing deterministic local calibration with a least-squares solver on the cantilever beam problem.

```

1 strategy
   single_method
3   tabular_graphics_data
   tabular_graphics_file 'cantilever_calibration.dat'
5
7 method
   nl2sol
9   convergence_tolerance 1.0e-6
#  output verbose
11
model
13  single
15 variables
   active design
17  continuous_design 3
   upper_bounds 31000000 10 10
19   initial_point 29000000 4 4
   lower_bounds 27000000 1 1
21   descriptors 'E' 'w' 't'
   continuous_state 3
23   initial_state 40000 500 1000
   descriptors 'R' 'X' 'Y'
25
interface
27  direct
   analysis_driver = 'mod_cantilever'
29
responses
31  calibration_terms 3
   calibration_data_file = 'dakota_cantilever_examples.clean.dat'
33   #calibration_data_file = 'dakota_cantilever_examples.error.dat'
   freeform
35  descriptors = 'area' 'stress' 'displacement'
   analytic_gradients
37  no_hessians

```

```

<<<<< Function evaluation summary: 18 total (17 new, 1 duplicate)
<<<<< Best parameters           =
                                2.8499999995e+07 E
                                2.5000000002e+00 w
                                2.9999999997e+00 t
                                4.0000000000e+04 R
                                5.0000000000e+02 X
                                1.0000000000e+03 Y
<<<<< Best residual norm =  2.9442886479e-06; 0.5 * norm^2 =  4.3344178210e-12
<<<<< Best residual terms      =
                                -2.3943602656e-10
                                2.9442885534e-06
                                7.0639355476e-10
<<<<< Best data captured at function evaluation 17

Confidence Interval for E is [  2.8499511022e+07,  2.8500488969e+07 ]
Confidence Interval for w is [  2.4999189436e+00,  2.5000810568e+00 ]
Confidence Interval for t is [  2.9999176961e+00,  3.0000823033e+00 ]

```

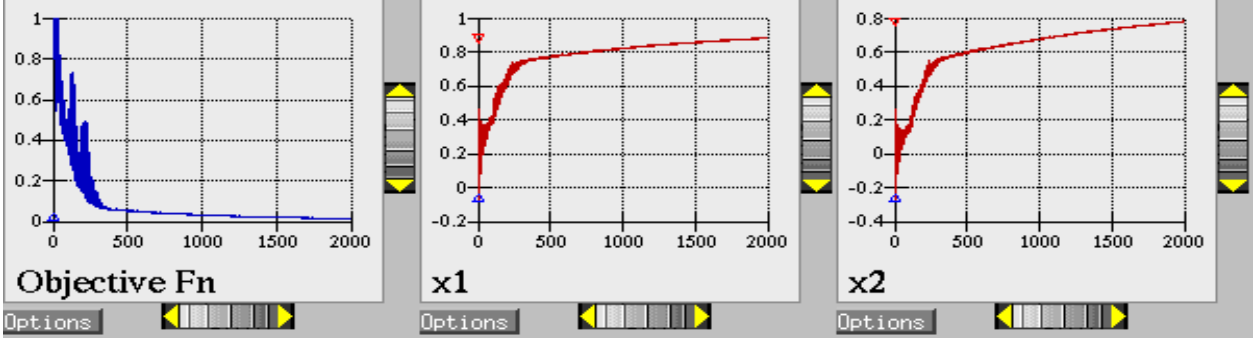
Figure 5.4: Dakota output showing optimal local gradient-based calibration result for cantilever beam, clean data.

evaluations for local methods, depending on the number of variables, and may require from thousands to tens-of-thousands of function evaluations for global methods. Given the computational cost, it is often prudent to use derivative-free methods to identify regions of interest and then use gradient-based methods to hone in on the solution. In addition to slow convergence, nonlinear constraint support in derivative-free methods is an open area of research and, while supported by many methods in Dakota, is not as refined as constraint support in gradient-based methods.

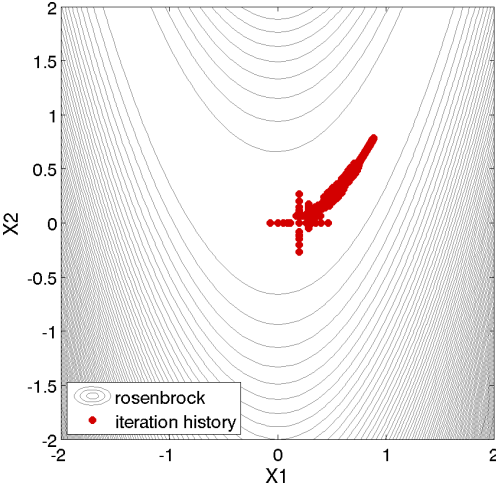
Local pattern search algorithms work by sampling the objective function at points on a stencil, often chosen to align with the coordinate axes. The stencil moves, expands, and contracts as the algorithm progresses. A typical pattern search iteration history is provided in Figures 5.5(a) and (b), which show the locations of the function evaluations used in the pattern search algorithm. Figure 5.5(c) provides a close-up view of the pattern search function evaluations used at the start of the algorithm (from a starting point $(x_1, x_2) = (0.0, 0.0)$). The coordinate pattern is clearly visible at the start of the iteration history, and the decreasing size of the coordinate pattern is evident as the design points move toward $(x_1, x_2) = (1.0, 1.0)$.

While pattern search algorithms are useful in many optimization problems, this example shows some of the drawbacks to this algorithm. While a pattern search method may make good initial progress towards an optimum, it is often slow to converge, here not fully converging after 2000 iterations. On a smooth, differentiable function such as that depicted, a nongradient-based method will not be as efficient as a gradient-based method. However, there are many engineering design applications where gradient information is inaccurate or unavailable, which renders gradient-based optimizers ineffective. Thus, pattern search algorithms are often good choices in complex engineering applications when the quality of gradient data is suspect.

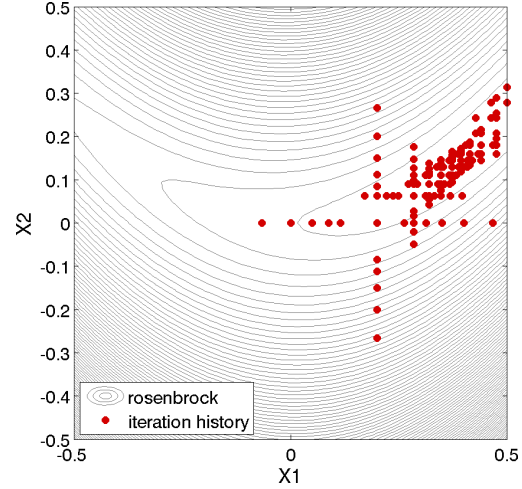
Recommended method: We recommend the Coliny Pattern Search algorithm for a derivative-



(a)



(b)



(c)

Figure 5.5: Pattern search optimization example: (a) screen capture of the Dakota graphics, (b) sequence of design points (dots) evaluated and (c) close-up view illustrating the shape of the coordinate pattern used.

free local method. A Dakota input file shown in Listing 5.3 applies a pattern search method to the cantilever beam optimization problem to minimize the objective function (area), subject to constraints on stress and displacement (see Section 2.1). The input is similar to the input file for the gradient-based optimization, except it has a different set of keywords in the method block of the input file (line 6–15), and the gradient specification in the responses block has been changed to `no_gradients` (line 39). The pattern search optimization algorithm used, `coliny_pattern_search` is part of the SCOLIB library [17]. See the Dakota Reference Manual [2] for more information on the *method* block commands that can be used with SCOLIB algorithms.

The tailing portion of the Dakota output shows the final design point and constraint values; see Figure 5.6. A solution similar to the gradient-based approach was found, though the constraint for displacement was slightly violated (this can often be addressed through larger constraint penalties in derivative-free methods; see commented option in the input file). However, the algorithm required about an order of magnitude more function evaluations.

Listing 5.3: Dakota input file showing local derivative-free optimization on the cantilever beam problem.

```

strategy
2   single_method
   tabular_graphics_data
4   tabular_graphics_file 'cantilever_opt_ps.dat'

6 method
   max_iterations = 1000
8   max_function_evaluations = 2000
   coliny_pattern_search
10  solution_accuracy = 1e-4
   initial_delta = 0.5
12  threshold_delta = 1e-4
   exploratory_moves basic_pattern
14  contraction_factor = 0.75
#   constraint_penalty = 10000
16

18 model
   single

20 variables
   continuous_design = 2
22   upper_bounds    4.0    4.0
   initial_point    2.5    2.5
24   lower_bounds    1.0    1.0
   descriptors      'w'    't'
26   continuous_state = 4
   initial_state    40000. 29.E+6 500. 1000.
28   descriptors      'R'    'E'    'X'    'Y'

30 interface
   direct
32   analysis_driver = 'mod_cantilever'

34 responses
   objective_functions = 1
36   # constraints assumed <= 0 unless bounds given
   nonlinear_inequality_constraints = 2
38   descriptors = 'area' 'stress' 'displacement'
   no_gradients
40   no_hessians

```

```

<<<<< Function evaluation summary: 149 total (149 new, 0 duplicate)
<<<<< Best parameters          =
                2.3875000000e+00 w
                3.2493181357e+00 t
                4.0000000000e+04 R
                2.9000000000e+07 E
                5.0000000000e+02 X
                1.0000000000e+03 Y
<<<<< Best objective function =
                7.7577470490e+00
<<<<< Best constraint values   =
                -1.6832688252e-01
                4.1748455156e-02
<<<<< Best data captured at function evaluation 142

```

Figure 5.6: Dakota output showing results from local derivative-free optimization with a pattern search algorithm.

5.2.3 Derivative-Free Global Methods

Dakota has a number of global optimization algorithms, including DiRECT, genetic algorithms, and surrogate-based approaches. Here we discuss genetic algorithms. In contrast to pattern search algorithms, which are local optimization methods, evolutionary algorithms (EA) are global optimization methods. EAs are best suited to optimization problems that have multiple local optima, and where gradients are either too expensive to compute or are not readily available.

Evolutionary algorithms work by generating an initial random sample in the parameter space, computing function values/constraints at those points, and then determining the next places to sample based on biological genetic selection principles of mutation and fitness. A simplified evolutionary or genetic algorithm for optimization would include:

1. Initialize a random population (sample) of individual x values.
2. Evaluate the fitness with respect to optimization objective and constraints.
3. Select more fit individuals as parents to reproduce.
4. Recombine and mutate to create a new population.
5. Iterate to (2.) until convergence is reached.

Here the implementation details of each step, for example how to select parents, how to combine and mutate, vary greatly among implementations, but are usually biologically inspired. Figure 5.7(a) shows the population of 50 randomly selected design points that comprise the first generation of the EA, and Figure 5.7(b) shows the final population of 50 design points, where most of the 50 points are clustered near $(x_1, x_2) = (0.98, 0.95)$.

Listing 5.4 shows a Dakota input file that uses an EA to solve the cantilever beam optimization problem as described in Section 2.1. Each generation of the EA has a population size of 50 (line 11).

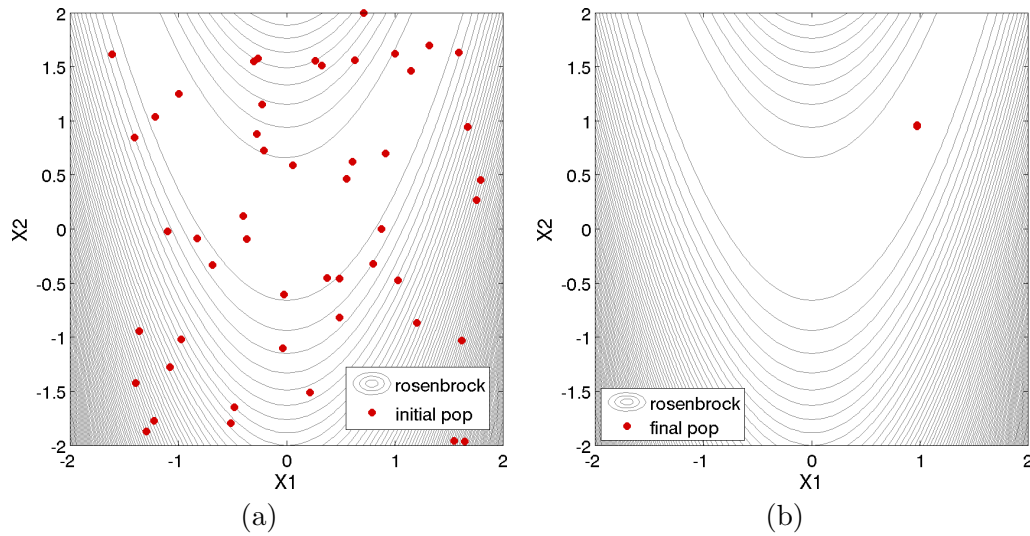


Figure 5.7: Evolutionary algorithm optimization example: 50 design points in the (a) initial and (b) final populations selected by the evolutionary algorithm.

The algorithm will take at most 100 iterations comprising no more than 2000 function evaluations (lines 7 and 8). The EA software available in Dakota provides the user with much flexibility in choosing the settings used in the optimization process. The details of all the settings are not discussed here; see [2] and [17].

On completion, the file `cantilever_opt_ea.dat` provides a listing of the design parameter values and objective function values for all 2,000 design points evaluated during the running of the EA. The final solution is shown in Figure 5.8.

```

<<<<< Function evaluation summary: 2007 total (2007 new, 0 duplicate)
<<<<< Best parameters      =
      2.3547960771e+00 w
      3.3249245714e+00 t
      4.0000000000e+04 R
      2.9000000000e+07 E
      5.0000000000e+02 X
      1.0000000000e+03 Y
<<<<< Best objective function =
      7.8295193371e+00
<<<<< Best constraint values =
      -6.8024421496e+02
      -3.4463344891e-03
<<<<< Best data captured at function evaluation 1982

```

Figure 5.8: Dakota output showing optimal point found by a global evolutionary algorithm.

EAs are better suited to optimization problems where conventional gradient-based optimization fails, such as situations where there are multiple local optima and/or gradients are not available. In

Listing 5.4: Dakota input file showing global optimization on the cantilever beam problem with an evolutionary algorithm.

```

strategy
2   single_method
   tabular_graphics_data
4   tabular_graphics_file 'cantilever_opt_ea.dat'

6 method
   max_iterations = 100
8   max_function_evaluations = 2000
   coliny_ea
10  seed = 11011011
   population_size = 50
12  fitness_type merit_function
   mutation_type offset_normal
14  mutation_rate 1.0
   crossover_type two_point
16  crossover_rate 0.0
   replacement_type chc = 10
18

model
20 single

22 variables
   continuous_design = 2
24   upper_bounds 4.0 4.0
   initial_point 2.5 2.5
26   lower_bounds 1.0 1.0
   descriptors 'w' 't'
28   continuous_state = 4
   initial_state 40000. 29.E+6 500. 1000.
30   descriptors 'R' 'E' 'X' 'Y'

32 interface
   direct
34   analysis_driver = 'mod_cantilever'

36 responses
   objective_functions = 1
38   # constraints assumed <= 0 unless bounds given
   nonlinear_inequality_constraints = 2
40   descriptors = 'area' 'stress' 'displacement'
   no_gradients
42   no_hessians

```

such cases, the computational expense of an EA is warranted since other optimization methods are not applicable or impractical. In many optimization problems, EAs often quickly identify promising regions of the design space where the global minimum may be located. However, an EA can be slow to converge to the optimum. For this reason, it can be an effective approach to combine the global search capabilities of an EA with the efficient local search of a gradient-based algorithm in a *hybrid optimization* strategy. In this approach, the optimization starts by using a few iterations of an EA to provide the initial search for a good region of the parameter space (low objective function and/or feasible constraints), and then it switches to a gradient-based algorithm (using the best design point found by the EA as its starting point) to perform an efficient local search for an optimum design point. More information on this hybrid approach is provided in the “Strategy” chapter of the Dakota User’s Manual [1].

Another effective method for global optimization, especially for costlier computational models, is the Efficient Global Optimization (EGO) approach. This is discussed in the section “Efficient Global Minimization” of the Dakota User’s Manual [1].

5.3 Summary and Additional Approaches

In selecting an optimization method, important considerations include the type of variables in the problem (continuous, discrete, mixed), whether a global search is needed or a local search is sufficient, and the required constraint support (unconstrained, bound constrained, or generally constrained). Less obvious, but equally important, considerations include the efficiency of convergence to an optimum (i.e., convergence rate) and the robustness of the method in the presence of challenging design space features (e.g., nonsmoothness).

Table 5.2 provides a more extensive reference for selecting from among all of Dakota’s optimization methods. Here blank fields inherit the values from above. With respect to constraint support, the methods with more advanced constraint support are also applicable to the lower constraint support levels; they are listed only at their highest level of constraint support for brevity. For example, all methods listed as supporting nonlinear constraints also support bound constraints.

Because of the computational cost of running simulation models, surrogate-based optimization (SBO) methods are often used to reduce the number of actual simulation runs. In SBO, a surrogate or approximate model is automatically constructed by Dakota based on a limited number of simulation runs. The optimization is then performed on the surrogate model. Dakota has an extensive framework for managing a variety of local, multipoint, global, and hierarchical surrogates for use in optimization. Finally, sometimes there are multiple objectives that one may want to optimize simultaneously instead of a single scalar objective. In this case, one may employ multi-objective methods described in “Optimization Capabilities” in the Dakota User’s Manual to either form a single composite objective, or assess the trade-off between multiple objectives directly.

Table 5.2: Detailed guidelines for selecting from among all Dakota optimization methods. Blank fields inherit the values from above.

Method Classification	Desired Problem Characteristics	Applicable Methods
Gradient-Based Local	smooth; continuous variables no constraints	optpp_cg
	smooth; continuous variables; bound constraints	dot_bfgs, dot_frcg, conmin_frcg
	smooth; continuous variables; bound constraints, linear and nonlinear constraints	npsol_sqp, nlpql_sqp, dot_mmfd, dot_slp, dot_sqp, conmin_mfd, optpp_newton, optpp_q_newton, optpp_fd_newton, weighted sums (multiobjective), pareto_set strategy (multiobjective)
Gradient-Based Global	smooth; continuous variables; bound constraints, linear and nonlinear constraints	hybrid_strategy, multi_start strategy
Derivative-Free Local	nonsmooth; continuous variables; bound constraints	optpp_pds
	nonsmooth; continuous variables; bound constraints, linear and nonlinear constraints	asynch_pattern_search, coliny_cobyala, coliny_pattern_search, coliny_solis_wets, surrogate_based_local
	nonsmooth; continuous variables; discrete variables; bound constraints, nonlinear constraints	mesh_adaptive_search
Derivative-Free Global	nonsmooth; continuous variables; bound constraints	ncsu_direct
	nonsmooth; continuous variables; bound constraints, linear and nonlinear constraints	coliny_direct, efficient_global, surrogate_based_global
	nonsmooth; continuous variables, discrete variables; bound constraints, linear and nonlinear constraints	coliny_ea, sogas, moga (multiobjective)

Chapter 6

Uncertainty Quantification

At a high level, uncertainty quantification (UQ) constitutes the process of characterizing input, numerical, and experimental uncertainties – consisting of both measurement errors and variability in replicate data, propagating these uncertainties through a computational model, and performing statistical or interval assessments on the resulting responses. This process determines the effect of uncertainties and assumptions on model responses or quantities of interest (QoI). In Section 6.1, we summarize techniques to propagate input uncertainties through models whereas in Section 6.2, we discuss Bayesian techniques to quantify input uncertainties.

For this discussion, inputs collectively refer to model parameters, initial conditions, boundary conditions, or exogenous forces. For models in which inputs are derived from closure, constitutive or phenomenological relations, one must employ model calibration techniques to estimate means, moments, or ultimately probability density functions (PDF) for these calibrated inputs based on experimental data or high-fidelity codes. This is often termed *inverse uncertainty quantification*. For CASL applications, inputs requiring calibration include cross-section values, closure and phenomenological parameters, and initial and boundary conditions.

The quantification of response or output uncertainties facilitates optimal design and decision making and is necessary to ensure robustness, performance or safety margins. For example, outputs specified for emergency core cooling systems include peak clad temperatures and maximum local cladding oxidation. The manner in which output statistics are employed depends on the application. For the assessment of design margins, the sample mean $\bar{x}$ and sample variance s^2 can be used to construct 2σ confidence intervals $\bar{x} \pm 2s$ whereas a predictive distribution for the output might be compared to replicate data to assess validation.

Wilks' formula constitutes a classical approach for assessing output uncertainty [29, 41]. In this coverage approach, the code is run N times for randomly selected input values chosen from expert-specified intervals. The outputs are then ranked to establish tolerance bounds for the response. This approach has the advantage that the number of required code evaluations is independent of the number of parameters. However, the resulting tolerance bounds can be overly conservative and the techniques detailed in Sections 6.1 and 6.2 yield more precise bounds or densities for inputs and outputs.

UQ is related to the sensitivity analysis, detailed in Chapter 3, in that the common goal is to gain an understanding of how variations in the parameters or inputs affect the response functions of the engineering design problem. However, for UQ, some or all of the components of the parameter vector are considered to be uncertain as specified by particular probability distributions (e.g.,

normal, exponential, extreme value), or other uncertainty structures. By assigning specific distributional structure to the inputs, distributional structure for the outputs (i.e., response statistics) can be inferred. UQ can thus be defined as the process of quantifying the imprecision of computed model responses or quantities of interest whereas sensitivity analysis ascertains how uncertainty in model outputs can be apportioned to uncertainties in model inputs when taken either singly or in combination over the range of input values.

6.1 Uncertainty Propagation

Whereas Dakota provides a number of options for propagating uncertainties through models, we focus on sampling and stochastic polynomial methods. The selection of these techniques to employ can be based on the following criteria:

- Sampling methods are applicable for nonsmooth and/or multi-modal response functions and general input densities including those for correlated parameters, initial conditions, boundary conditions or exogenous forces. Since sampling methods require numerous model evaluations – e.g., hundreds to millions – they necessitate that models be efficient to evaluate or require the use of suitable surrogate models. Sampling methods yield response samples, which can be post-processed – e.g., using kernel density estimation (kde) routines – to construct response densities.
- Stochastic polynomial methods require smooth response functions. Stochastic collocation (SC) methods are applicable for general input densities whereas nonintrusive polynomial chaos expansions (PCE) require the specification or construction of orthogonal polynomials. For normal or uniform input densities, Hermite or Legendre polynomials are employed. For densities that do not correspond to members in the Askey family of polynomials, Dakota provides the capability of employing empirical histograms to generate orthogonal polynomials. Evaluation of the quantity of interest often necessitates that inputs are mutually independent, which is generally not the case, but can often be achieved through Nataf transformations. These methods utilize Dakota sparse grid routines to provide highly efficient evaluation of response moments and Sobol global sensitivity indices.

We note that Dakota provides a number of other techniques – including local and global reliability methods, interval methods, and mixed UQ algorithms – to propagate uncertainties through models. Readers are referred to the “Uncertainty Quantification Capabilities” chapter of [1] for details about these methods.

We do not differentiate between aleatoric uncertainties, which are inherent to a problem or experiments and are intrinsically probabilistic in nature, and epistemic uncertainties, which are due to lack of knowledge. We refer readers to [34] for details regarding the nature of these uncertainties and to the “Uncertainty Quantification Capabilities” chapter of [1] for a description of how Dakota algorithms accommodate these two classes of uncertainties. Further details regarding various uncertainty propagation techniques are provided in [34, 37].

6.1.1 Sampling Methods

Sampling-based methods are the most robust uncertainty techniques available, are applicable to almost all simulations, and possess rigorous error bounds. Consequently, they should be used when-

ever the function is relatively inexpensive to compute and adequate sampling can be performed. In the case of computationally expensive simulations, however, the number of function evaluations required by traditional techniques such as Monte Carlo and Latin hypercube sampling (LHS) quickly becomes prohibitive, especially if tail statistics are needed. We note that the issues associated with tail statistics can be mitigated through the use of importance sampling.

Alternatively, one can apply the traditional sampling techniques to a surrogate function approximating the expensive computational simulation. However, if this approach is selected, the user should be aware that it is difficult to assess the accuracy of the results. Unlike the case of surrogate-based local minimization, there is no simple pointwise calculation to verify the accuracy of the approximate results. This is due to the functional nature of uncertainty quantification; i.e., the accuracy of the surrogate over the entire parameter space needs to be considered, not just around a candidate optimum as in the case of surrogate-based local optimization. This issue especially manifests itself when trying to estimate low probability events such as the catastrophic failure of a system.

Due to the computational complexity of CASL codes, sampling methods will generally need to be applied to surrogate models rather than physics-based codes. Hence the issues associated with establishing the accuracy of the surrogate must be addressed or the accuracy of the surrogate-based sampling results verified for these codes. Techniques to establish surrogate accuracy include “leave-one-out” cross-validation for Gaussian processes [6] and the Dakota K -fold cross-validation capability, which includes “leave-one-out.” Details regarding the use of the Dakota sampling capability are provided in Section 6.1.5.

6.1.2 Stochastic Polynomial Methods

Stochastic polynomial methods comprise a second class of forward propagation techniques, which are available in Dakota. The development of these techniques mirrors that of deterministic Galerkin and finite element analysis utilizing the notions of projection, collocation, orthogonality, and weak convergence [12, 13]. Rather than providing point estimates, they form an approximation to the functional relationship between random inputs and response functions, which provides a representation of output uncertainties for multi-code simulations. Expansion methods include nonintrusive polynomial chaos expansions (PCE), which employ multivariate orthogonal polynomials that are tailored to particular input probability distributions, and stochastic collocation (SC), which employs multivariate interpolation polynomials.

For certain applications, sampling-based models can be efficiently combined with stochastic polynomial methods. For example, it is often advantageous to employ Latin hypercube sampling to construct a surrogate, which is subsequently employed with a regression-based nonintrusive polynomial chaos expansion.

Nonintrusive Polynomial Chaos Expansions (PCE)

To motivate nonintrusive PCE methods, which for certain implementation regimes are also termed pseudo-spectral or discrete projection methods, we consider a parameter-dependent response $Y(X)$ where $X \in \mathbb{R}^M$. For random inputs X , this response is represented by the truncated expansion

$$Y(X) = \sum_{j=0}^J \alpha_j \Psi_j(X) \quad (6.1)$$

where $\Psi_j(X)$ are polynomials that are orthogonal with respect to inner products corresponding to common probably density functions. For example, Hermite and Legendre polynomials with weights $e^{-x^2/2}$ and 1 are respectively used to represent single-variate normal and uniform distributions. As detailed in the “Stochastic Expansion Methods” chapter of the Dakota Theory Manual [3] and Chapter 10 of [34], tensored polynomials are constructed as basis functions for multivariate densities.

If we denote the density by $\rho(x)$ and note that $\Psi_0(X) = 1$, it follows that

$$E[\Psi_0(X)] = 1$$

and

$$\begin{aligned} E[\Psi_i(X)\Psi_j(X)] &= \int_{\Gamma} \Psi_i(x)\Psi_j(x)\rho(x)dx \\ &= \delta_{ij}\gamma_i \end{aligned}$$

where $\Gamma = [0, 1]^M$ for scaled Legendre polynomials, δ_{ij} is the Kronecker delta and the normalization factor is

$$\gamma_i = E[\Psi_i^2(X)].$$

We note that γ_i can be computed analytically for each polynomial in the Askey family, which includes Hermite and Legendre polynomials.

Based on these orthogonality properties, it follows that the mean and variance of Y are

$$\begin{aligned} E[Y(X)] &= \alpha_0 \\ \text{Var}[Y(X)] &= \sum_{j=1}^J \alpha_j^2 \gamma_j. \end{aligned} \tag{6.2}$$

Hence these values, as well as higher order moments, can be computed very efficiently once one has constructed the coefficients α_j .

For the nonintrusive PCE method, one takes the weighted inner product of (6.1) with respect to Ψ_j and enforces orthogonality to obtain

$$\alpha_j = \frac{1}{\gamma_j} \int_{\Gamma} Y(x)\Psi_j(x)\rho(x)dx.$$

Hence the determination of the coefficients α_j requires numerical quadrature over $\Gamma \subset \mathbb{R}^M$. In Dakota, this is achieved using tensored Gaussian or sparse grid quadrature techniques. The evaluation of

$$\alpha_j \approx \frac{1}{\gamma_j} \sum_{i=1}^{N_q} Y(x_i)\Psi_j(x_i)\rho(x_i)w_i \tag{6.3}$$

thus employs codes nonintrusively, or as a black box, to evaluate the response at parameter values x_i . Details regarding this method can be found in Chapter 10 of [34]. Alternatively, one can determine coefficients using Dakota’s linear regression capabilities as detailed in the “Stochastic Expansion Methods” chapter of the Dakota Theory Manual [3]. The implementation of this technique in Dakota is illustrated in Section 6.1.5.

Stochastic Collocation (SC)

In the stochastic collocation method, one represents the response for random inputs X as

$$Y(X) = \sum_{j=1}^J r_j L_j(X)$$

where $r_j = Y(x_j)$ is the response value at the interpolation point x_j and $L_j(x)$ is a Lagrange polynomial. In 1-D, the Lagrange polynomial can be represented as

$$L_j(x) = \prod_{\substack{k=1 \\ k \neq j}}^J \frac{x - x_k}{x_j - x_k}$$

which highlights the property that $L_j(x_i) = \delta_{ij}$. For moderate parameter dimensionality, Dakota provides the capability for implementing multivariate interpolation on Smolyak sparse grids. As detailed in the “Stochastic Expansion Methods” chapter of the Dakota Theory Manual [3], Dakota also provides the capability for implementing local or global interpolating polynomials and either value-based or gradient-enhanced representations. The implementation of stochastic collocation in Dakota is illustrated in Section 6.1.5.

Table 6.1 provides a reference for choosing a Dakota propagation method or strategy based on the properties of the model.

Table 6.1: Guidelines for uncertainty propagation method selection.

Method Classification	Desired Problem Characteristics	Applicable Methods
Sampling	Nonsmooth and/or multi-modal response functions; general densities; computationally efficient models or surrogates	sampling
Stochastic Polynomial	Smooth response functions; can be combined with Dakota sparse grid routines	polynomial_chaos stoch_collocation

6.1.3 Verification

Verification comprises a critical component of uncertainty propagation when quantifying response variability. For many problems, one or more of the following steps can be employed to verify response uncertainties. The use of this verification strategy is illustrated in Section 6.1.5.

- (i) Compare the mean, variance, skewness and kurtosis provided by noninvasive PCE and stochastic collocation.
- (ii) Construct the response densities using the sampling methods discussed in Section 6.1.1 and compare the resulting moments with those computed in (i).

- (iii) For linearly parameterized problems with uncorrelated Gaussian inputs $X_i \sim \mathcal{N}(\mu_i, \sigma_i^2)$, the response $Y = \sum_{i=1}^N a_i X_i$ is normally distributed with mean and variance

$$E(Y) = \sum_{i=1}^N a_i \mu_i, \quad \text{var}(Y) = \sum_{i=1}^N a_i^2 \sigma_i^2. \quad (6.4)$$

These moments can be compared with those constructed in (i) and (ii).

6.1.4 Prediction Intervals

One goal when propagating uncertainties is to construct credible or prediction intervals for the model response or quantity of interest. This can be achieved using the methods of Section 6.1.1 to sample from input densities, constructed either experimentally or using the Bayesian model calibration techniques detailed in Section 6.2. The propagation of solely input density information yields credible intervals, which quantify the accuracy of the model. The simultaneous propagation of input and experimental uncertainties – e.g., using the statistical model (6.5) – yields prediction intervals, which quantify the probability of observing the next experimental measurement or numerical simulation. For example, the interval (a, b) is a 95% prediction interval if the probability that a future experimental observation y_i falls within (a, b) is at least 0.95. For this reason, prediction intervals are typically preferable to credible intervals when experimentally validating model behavior. Details regarding posterior predictive distributions can be found in [11].

6.1.5 Uncertainty Propagation: Cantilever Beam Example

We illustrate the sampling and stochastic polynomial methods for the cantilever beam example detailed in Section 2.1. Specifically, we illustrate the use of random sampling, noninvasive polynomial chaos expansions, and stochastic collocation to propagate parameter uncertainties through the model to quantify uncertainties in the output area, stress, and displacement. In all cases, parameters were assumed to have the normal distributions $R \sim \mathcal{N}(4.0E5, 4.0E6)$, $E \sim \mathcal{N}(2.9E7, 2.1025E12)$, $X \sim \mathcal{N}(5.0E2, 1.0E4)$ and $Y \sim \mathcal{N}(1.0E3, 1.0E4)$. Additional details regarding the performance of random sampling, PCE, and stochastic collocation, for this example, are provided in [37].

Random Sampling

We first illustrate the use of random sampling to construct densities for the output based on 10^6 samples from the assumed normal input distributions. The input file is shown in Listing 6.1.

An excerpt of the random sampling output is shown in Figure 6.1. First the mean, standard deviation, skewness, and kurtosis moments are calculated for each of the three response functions, as well as 95% confidence intervals for each response mean and standard deviation. Correlation matrices among all inputs and outputs are provided to determine the degree of linear relationships among variables. Finally, the full output provides a listing of CDF probabilities.

Nonintrusive Polynomial Chaos Expansions (PCE)

A typical Dakota input file for performing uncertainty propagation using nonintrusive PCE is shown in Listing 6.2. In this example, we compute CDF probabilities for 17 probability levels of the cantilever beam equations. Due to the low parameter dimensionality, we select Gaussian

Listing 6.1: UQ input for random sampling.

```

strategy
2   tabular_graphics_data
   tabular_graphics_file = 'cantilever_sampling.dat'
4   single_method

6 method
   sampling
8   sample_type random
   samples = 1000000
10  seed = 17
   response_levels = 1. 5. 10.
12  10000. 20000. 40000.
   1. 2. 3.
14

model
16  single

18 variables
   active uncertain
20  continuous_design = 2
   initial_point    2.5    2.5
22  descriptors      'w'    't'
   normal_uncertain = 4
24  means            = 40000. 29.E+6 500. 1000.
   std_deviations    = 2000.  1.45E+6 100. 100.
26  descriptors      = 'R'    'E'    'X'  'Y'

28 interface
   analysis_driver = 'mod_cantilever'
30  direct

32 responses
   response_functions = 3
34  descriptors = 'area' 'stress' 'displacement'
   no_gradients
36  no_hessians

```


Statistics based on 1000 samples:

Moment-based statistics for each response function:

	Mean	Std Dev	Skewness	Kurtosis
area	6.2500000000e+00	0.0000000000e+00	0.0000000000e+00	-3.0000000000e+00
stress	1.7353650406e+04	5.8085356413e+03	-6.5053618765e-03	7.3520733236e-02
displacement	1.7161216049e+00	4.1203461736e-01	1.3498700788e-01	-1.6987570074e-02

95% confidence intervals for each response function:

	LowerCI_Mean	UpperCI_Mean	LowerCI_StdDev	UpperCI_StdDev
area	6.2500000000e+00	6.2500000000e+00	0.0000000000e+00	0.0000000000e+00
stress	1.6993203553e+04	1.7714097260e+04	5.5646488747e+03	6.0749462045e+03
displacement	1.6905529262e+00	1.7416902836e+00	3.9473425169e-01	4.3093273234e-01

Simple Correlation Matrix among all inputs and outputs:

	R	E	X	Y	area	stress	displacement
R	1.00000e+00						
E	-1.09891e-03	1.00000e+00					
X	-1.89011e-03	3.84943e-04	1.00000e+00				
Y	2.28006e-03	1.47231e-03	-1.28899e-04	1.00000e+00			
area	0.00000e+00	0.00000e+00	0.00000e+00	0.00000e+00	1.00000e+00		
stress	-3.46903e-01	1.62695e-03	6.50987e-01	6.74942e-01	0.00000e+00	1.00000e+00	
disp.	1.99216e-03	-4.87018e-01	3.80836e-01	7.81233e-01	0.00000e+00	7.74982e-01	1.00000e+00

Partial Correlation Matrix between input and output:

	area	stress	displacement
R	-0.00000e+00	-1.00000e+00	5.02670e-03
E	-0.00000e+00	-1.67347e-09	-9.87430e-01
X	-0.00000e+00	1.00000e+00	9.79611e-01
Y	-0.00000e+00	1.00000e+00	9.95042e-01

Simple Rank Correlation Matrix among all inputs and outputs:

	R	E	X	Y	area	stress	displacement
R	1.00000e+00						
E	1.29417e-02	1.00000e+00					
X	-1.02171e-02	3.68951e-03	1.00000e+00				
Y	3.49896e-03	5.36687e-03	4.26256e-03	1.00000e+00			
area	4.84427e-02	3.11778e-02	-2.69348e-02	-4.07852e-02	1.00000e+00		
stress	-3.31948e-01	-7.20001e-08	6.39440e-01	6.52457e-01	-6.90457e-02	1.00000e+00	
disp.	-1.33517e-03	-4.74628e-01	3.68797e-01	7.57213e-01	-5.19435e-02	7.57396e-01	1.00000e+00

Partial Rank Correlation Matrix between input and output:

	area	stress	displacement
R	4.79912e-02	-7.97285e-01	2.39053e-02
E	3.09448e-02	-6.40293e-03	-8.86932e-01
X	-2.64483e-02	9.31104e-01	8.26712e-01
Y	-4.10868e-02	9.34429e-01	9.49696e-01

Figure 6.1: Excerpt of UQ output for random sampling.

quadrature with a total of 57 function evaluations to compute the nonintrusive PCE coefficients (6.3), using tensor product quadrature points. The tensor product generates all combinations of values from each individual dimension so it is an all-way pairing of points. We note that one

would replace tensored Gaussian quadrature with sparse grid techniques for moderate parameter dimensionality; e.g., $M = 5$ to approximately $M = 40$, depending on the regularity of the modeling equations.

Once the expansion coefficients have been calculated, some statistics can be computed analytically – e.g., via (6.2) – whereas others must be evaluated numerically. For the numerical portion, the input file specifies the use of 10000 samples, which will be evaluated on the expansion to compute the CDF probabilities. We summarize in Figure 6.2 excerpts of the output. The full output lists a summary of the PCE coefficients, which reproduce the function for a Hermite polynomial basis. The analytic statistics for mean, standard deviation, and covariance are then presented for each of the response functions: area, stress, and displacement. Finally, we note the numerical results for the CDF probabilities based on 10^5 samples performed on the expansion. For example, approximately 50% of the displacement samples are determined to be less than or equal to 1.709 inches.

Stochastic Collocation

Here we illustrate the use of stochastic collocation built on an anisotropic sparse grid defined from numerically-generated orthogonal polynomials; see Chapters 10 and 11 of [34] for details regarding stochastic collocation. The input file is shown in Listing 6.3. In this example, we again compute CDF probabilities of stress and displacement for varying response levels. This example requires 233 function evaluations to compute the interpolating polynomials used for stochastic collocation.

Once the expansion coefficients have been calculated, some statistics are available analytically and others must be evaluated numerically. For the numerical portion, the input file specifies the use of 10000 samples, which will be evaluated on the expansion to compute the CDF probabilities. We summarize in Figure 6.3 excerpts from the output. We first note the moment statistics for mean, standard deviation, skewness, and kurtosis computed by numerical integration (see the “Analytic Moments” section in the “Stochastic Expansion Methods” chapter in the Dakota Theory Manual [3]), where the numerical row corresponds to integration using the original response values and the expansion row corresponds to integration using values from the interpolant. The response covariance and global sensitivity indices (Sobol indices) are presented next. This example shows that for stress, the variables R , X , and Y all play significant roles, but that the interactions between them are relatively negligible. For displacement, E , X , and Y have significant influence, while the interactions between E and X , E and Y , X and Y , and the three way interaction have significantly less effect. Finally, in the full output, we see the numerical results for the CDF probabilities based on 10000 samples performed on the expansion.

Method Verification

We illustrate here the verification techniques detailed in Section 6.1.3. We compile in Tables 6.2 and 6.3 the first four moments provided by random sampling, noninvasive PCE and stochastic collocation. Additionally, since stress is a linear function of the normally distributed inputs X , Y and R , the relations (6.2) can be used to compute analytic values for the stress moments.

We note that the moments for each of the three methods are in close agreement, with the exception of the third and fourth moments in the stress output. Since the stress is linearly dependent on the parameters, it is expected that the PCE and collocation methods should produce exact results for these two moments, neglecting rounding error. By comparison, the skewness and kurtosis,

Listing 6.2: UQ input for nonintrusive polynomial chaos expansions.

```

1 # Dakota Input File: cantilever_uq_pce.in
3 strategy
4   single_method
5   tabular_graphics_data
6     tabular_graphics_file = 'cantilever_uq_pce.dat'
7   graphics
9 method
10  polynomial_chaos
11    sparse_grid_level = 2 #non_nested
12    sample_type lhs
13    seed 12347
14    samples = 10000
15    num_probability_levels = 0 17 17
16    probability_levels =
17      .001 .01 .05 .1 .15 .2 .3 .4 .5 .6 .7 .8 .85 .9 .95 .99 .999
18      .001 .01 .05 .1 .15 .2 .3 .4 .5 .6 .7 .8 .85 .9 .95 .99 .999
19    cumulative_distribution
21 model
22   single
23 variables
24   active_uncertain
25   continuous_design = 2
26     initial_point 2.5 2.5
27     descriptors 'w' 't'
28   normal_uncertain = 4
29     means = 40000. 29.E+6 500. 1000.
30     std_deviations = 2000. 1.45E+6 100. 100.
31     descriptors = 'R' 'E' 'X' 'Y'
33 interface
34   direct
35     analysis_driver = 'mod_cantilever'
37 responses
38   response_functions = 3
39   descriptors = 'area' 'stress' 'displacement'
40   no_gradients
41   no_hessians

```

```

Statistics derived analytically from polynomial expansion:
Moment-based statistics for each response function:
                Mean                Std Dev                Skewness                Kurtosis
stress
  expansion:    1.7600000000e+04  5.7871581973e+03
  numerical:    1.7600000000e+04  5.7871581973e+03  9.4100742175e-15  6.2172489379e-15
displacement
  expansion:    1.7201243431e+00  4.0644795983e-01
  numerical:    1.7201243431e+00  4.0644787032e-01  1.5009952217e-01  4.9005496977e-02
Covariance matrix for response functions:
[[ 7.3191614098e-30  5.3786836096e-13  6.5282177146e-17
   5.3786836096e-13  3.3491200000e+07  1.8163897036e+03
   6.5282177146e-17  1.8163897036e+03  1.6519994405e-01 ]]

Cumulative Distribution Function (CDF) for stress:
  Response Level  Probability Level  Reliability Index  General Rel Index
  -----
  2.4921421856e+02  1.0000000000e-03
  4.1489075797e+03  1.0000000000e-02
  7.9708753041e+03  5.0000000000e-02
  1.0090342657e+04  1.0000000000e-01
  1.1589780322e+04  1.5000000000e-01
  1.2731567123e+04  2.0000000000e-01
  1.4564078343e+04  3.0000000000e-01
  1.6151010310e+04  4.0000000000e-01
  1.7689441098e+04  5.0000000000e-01
  1.9129203866e+04  6.0000000000e-01
  2.0683233939e+04  7.0000000000e-01
  2.2457356004e+04  8.0000000000e-01
  2.3589089220e+04  8.5000000000e-01
  2.4920875151e+04  9.0000000000e-01
  2.7044322788e+04  9.5000000000e-01
  3.0752664401e+04  9.9000000000e-01
  3.5331778223e+04  9.9900000000e-01

Cumulative Distribution Function (CDF) for displacement:
  Response Level  Probability Level  Reliability Index  General Rel Index
  -----
  5.8392829293e-01  1.0000000000e-03
  8.2796204947e-01  1.0000000000e-02
  1.0598405267e+00  5.0000000000e-02
  1.2097152423e+00  1.0000000000e-01
  1.2968601525e+00  1.5000000000e-01
  1.3747889667e+00  2.0000000000e-01
  1.4999245316e+00  3.0000000000e-01
  1.6076201989e+00  4.0000000000e-01
  1.7093348267e+00  5.0000000000e-01
  1.8118345186e+00  6.0000000000e-01
  1.9241773807e+00  7.0000000000e-01
  2.0596369386e+00  8.0000000000e-01
  2.1417150361e+00  8.5000000000e-01
  2.2453314112e+00  9.0000000000e-01
  2.3964502080e+00  9.5000000000e-01
  2.7290918315e+00  9.9000000000e-01
  3.0882954345e+00  9.9900000000e-01

```

Figure 6.2: Excerpt of UQ output for nonintrusive polynomial chaos expansion.

Listing 6.3: UQ input for stochastic collocation.

```

1 strategy
  single_method
3  tabular_graphics_data
    tabular_graphics_file = 'cantilever_uq_sc.dat'
5
method
7  stoch_collocation
    sparse_grid_level = 3
9    samples = 10000 seed = 12347 rng rnum2
    response_levels = 1. 5. 10. 15. 20. 50.
11   1. 10. 5000. 10000. 50000. 100000.
    .1 1. 2. 3. 5. 10.
13   variance_based_decomp
    output silent
15
variables
17  active_uncertain
    continuous_design = 2
19    initial_point    2.5    2.5
    descriptors      'w'    't'
21  normal_uncertain = 4
    means             = 40000. 29.E+6 500. 1000.
23    std_deviations  = 2000.  1.45E+6 100. 100.
    descriptors      = 'R'    'E'    'X'  'Y'
25
interface
27  direct
    analysis_driver = 'mod_cantilever'
29
responses
31  response_functions = 3
    descriptors = 'area' 'stress' 'displacement'
33    no_gradients
    no_hessians

```

Statistics derived analytically from polynomial expansion:

Moment-based statistics for each response function:

	Mean	Std Dev	Skewness	Kurtosis
stress				
expansion:	1.7600000000e+04	5.7871581973e+03	-1.3819173712e-14	5.3290705182e-15
numerical:	1.7600000000e+04	5.7871581973e+03	-1.8416777990e-14	5.7731597281e-15
displacement				
expansion:	1.7201241681e+00	4.0644979045e-01	1.5036497320e-01	6.8782204635e-02
numerical:	1.7201241681e+00	4.0644979045e-01	1.5036497320e-01	6.8782204635e-02

Covariance matrix for response functions:

```
[[ 7.8886090522e-31 -3.2350298255e-26 -6.5675038652e-31
  -3.2350298255e-26  3.3491200000e+07  1.8163966854e+03
  -6.5675038652e-31  1.8163966854e+03  1.6520143216e-01 ]]
```

Global sensitivity indices for each response function:

stress Sobol indices:

	Main	Total	
	1.1943435888e-01	1.1943435888e-01	R
	-4.9831897744e-14	1.3881742943e-13	E
	4.4028282056e-01	4.4028282056e-01	X
	4.4028282056e-01	4.4028282056e-01	Y
Interaction			
	-2.3136238238e-14		R E
	7.1188425348e-15		R X
	7.1188425348e-14		E X
	-8.1866689150e-14		R Y
	1.2635945499e-13		E Y
	-9.0765242319e-14		X Y
	1.6017395703e-14		R E X
	2.4915948872e-14		R E Y
	1.8508990590e-13		R X Y
	-2.6695659506e-14		E X Y

displacement Sobol indices:

	Main	Total	
	3.0376298835e-13	-1.3978473800e-13	R
	2.4260687756e-01	2.4452448915e-01	E
	1.5243149082e-01	1.5378097334e-01	X
	6.0208388530e-01	6.0457471845e-01	Y
Interaction			
	2.4193512347e-13		R E
	-2.5537596366e-13		R X
	3.8691316858e-04		E X
	-8.0376224352e-13		R Y
	1.5282637946e-03		E Y
	9.6013473508e-04		X Y
	-1.6666641839e-13		R E X
	1.9623626681e-13		R E Y
	3.4408550893e-13		R X Y
	2.4346209458e-06		E X Y

Figure 6.3: Excerpt of UQ output for stochastic collocation.

obtained with random sampling, are non-negligible but have values that are still five to seven orders of magnitude smaller than the sample mean and variance.

Stress				
	Mean	Std. Deviation	Skewness	Kurtosis
Analytic	1.760e+04	5.787e+03	–	–
Sampling	1.735e+04	5.809e+03	-6.505e-03	7.352e-02
PCE	1.760e+04	5.787e+03	9.410e-15	6.217e-15
Collocation	1.760e+04	5.787e+03	-1.842e-14	5.773e-15

Table 6.2: Comparison of moments for stress from random sampling, polynomial chaos expansion, and stochastic collocation methods.

Displacement				
	Mean	Std. Deviation	Skewness	Kurtosis
Sampling	1.716	4.120e-01	1.350e-01	-1.700e-02
PCE	1.720	4.064e-01	1.501e-01	4.901e-02
Collocation	1.720	4.064e-01	1.504e-01	6.878e-02

Table 6.3: Comparison of moments for displacement from random sampling, polynomial chaos expansion, and stochastic collocation methods.

The stress density constructed by random sampling is compared in Figure 6.4(a) with normal densities whose mean and variance (6.2) are constructed using coefficients α_j determined by nonintrusive PCE or stochastic collocation. The displacement density constructed by random sampling is plotted in Figure 6.4(b). We cannot compare to Gaussian representations with mean and variance computed using nonintrusive PCE or collocation coefficients since the nonlinear displacement dependence on inputs yields a non-Gaussian output as demonstrated by the magnitude of the skewness and kurtosis coefficients in Table 6.3. However, we could employ the nonintrusive PCE or collocation representations as surrogate models, which could subsequently be used with random sampling to more efficiently construct the non-Gaussian displacement density.

6.2 Bayesian Model Calibration

We noted in the introduction to this chapter that uncertainty quantification is broadly comprised of two steps: (i) quantification of uncertainties associated with models, inputs and experiments, and (ii) propagation of these uncertainties through models to quantify uncertainties in responses or quantities of interest. Here inputs refer to model parameters, initial conditions, boundary conditions or exogenous forces. We discuss here techniques to calibrate inputs derived from closure, constitutive or phenomenological relations for which values derived from fundamental principles are lacking.

The deterministic calibration techniques detailed in Chapter 5 provide point estimates for calibration inputs, but generally no measure of uncertainty. An exception is gradient-based methods, which produce asymptotic 95% confidence intervals for each input. The usefulness of these intervals

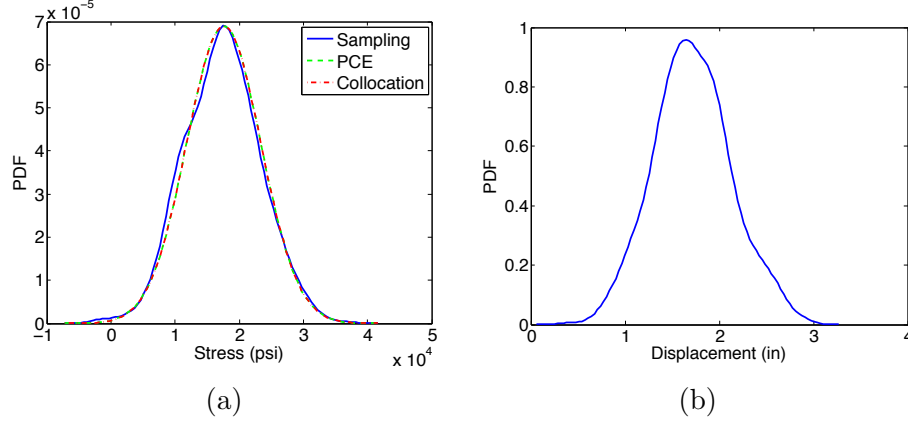


Figure 6.4: (a) Stress density constructed by random sampling and normal densities with mean and variance (6.2) computed using coefficients α_j determined by nonintrusive PCE and stochastic collocation. (b) Displacement density constructed by random sampling.

is tied to the validity of asymptotic assumptions applied to the problem at hand. As noted in the introduction to this chapter, approaches such as Wilks' formula employ uniform input distributions, which are generally based on expert opinion. These intervals are often based on qualitative, rather than quantitative, knowledge and hence they are typically conservative.

Bayesian inference provides a framework for probabilistic model calibration based on the assumption that calibration inputs are random variables having associated PDFs. These PDFs quantify both the support, or admissible parameter values, and the plausibility of each admissible parameter value. In Bayesian model calibration, one employs a likelihood, which incorporates measured data and computed model information, to update prior density information to obtain a more accurate posterior parameter density, which is consistent with experimental uncertainties.

Input densities or bounds, constructed in this manner, are tighter and contain more information than uniform densities constructed solely to bound potential input values. Propagation of these input densities using the sampling, nonintrusive polynomial chaos expansions, or stochastic collocation techniques of Section 6.1 will provide reduced response uncertainties and hence tighter robustness, performance or safety margins. For example, these densities could be employed in Wilks' formula to construct tighter tolerance bounds than those obtained using conservative, non-inference based input densities.

We summarize pertinent details required for implementation of the methods and refer readers to [20, 34] for additional examples and details regarding the theory and algorithms.

6.2.1 Direct Implementation of Bayes' Relation

To set notation, we consider a random calibration parameter vector $\Theta = [\Theta_1, \dots, \Theta_M]$ with the realization $\theta = [\theta_1, \dots, \theta_M]$. We consider the statistical model

$$D_i = y_i(\Theta) + \varepsilon_i, \quad i = 1, \dots, N \quad (6.5)$$

where $D = [D_1, \dots, D_N]$ denotes unobserved (random) data, $y(\Theta) = [y_1(\Theta), \dots, y_N(\Theta)]$ is the parameter-dependent model, and $\varepsilon = [\varepsilon_1, \dots, \varepsilon_N]$ is a random vector, which represents experimental and model errors. Throughout this discussion, we assume that ε_i are independently and iden-

tically distributed (iid) and $\varepsilon_i \sim \mathcal{N}(0, \sigma^2)$ where the experimental error variance σ^2 is determined experimentally or estimated through the inference procedure. We note that the model may additionally depend on spatial or temporal independent variables – e.g., $y_i = y(x_i, \Theta)$ or $y_i = y(t_i, \Theta)$ – but we simplify notation by suppressing these latter dependencies since model calibration focuses on uncertain parameters.

In Bayesian inference, one employs Bayes' relation

$$\pi(\theta|d) = \frac{\mathcal{L}(\theta; d)\pi_0(\theta)}{\int_{\mathbb{R}^M} \mathcal{L}(\theta; d)\pi_0(\theta)d\theta} \quad (6.6)$$

for observed data d to update a prior density $\pi_0(\theta)$, using the likelihood $\mathcal{L}(\theta; d)$, to obtain a more informative posterior density $\pi(\theta|d)$.

Prior Density. The prior density $\pi_0(\theta)$ incorporates any knowledge that one has about parameters prior to obtaining observations d . This could come from previous similar experiments or analysis regarding similar models. If prior knowledge is of questionable accuracy, it is better to use a noninformative prior, which is often taken as an improper uniform density posed on the parameter support. For example, one would employ $\pi_0(\theta) = \chi_{(0, \infty)}(\theta)$ for positive parameters, where $\chi_{(0, \infty)}(\theta)$ is the characteristic function having a value of 1 for $\theta \in (0, \infty)$ and 0 for $\theta \in (-\infty, 0]$. In analyses assuming random experimental error variance σ^2 , a standard noninformative prior for σ^2 is the Jeffreys prior $\pi_0(\sigma^2) \propto (1/\sigma^2)$.

Likelihood Function. The likelihood $\mathcal{L}(\theta; d)$ incorporates information provided by the samples and model and constitutes the mechanism through which data informs the posterior density. The likelihood can be interpreted as quantifying the probability of obtaining the observations d for a given value θ of the parameter Θ . The likelihood can generally be written as

$$\mathcal{L}(\theta; d) = f[y(\theta) - d]$$

where the function f can be constructed to emphasize specific relations between the model and data. For the statistical model (6.5) with iid errors ε_i that are normally distributed, $\varepsilon_i \sim \mathcal{N}(0, \sigma^2)$, one employs the likelihood relation

$$\begin{aligned} \mathcal{L}(\theta; d) &= \prod_{i=1}^N \frac{1}{\sigma\sqrt{2\pi}} e^{-(d_i - y_i(\theta))^2/2\sigma^2} \\ &= \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-SS_q/2\sigma^2} \end{aligned} \quad (6.7)$$

where

$$SS_q = \sum_{i=1}^N [d_i - y_i(\theta)]^2$$

denotes the sum of squares error. We note that the likelihood and prior density must be specified by users when employing sampling-based Bayesian algorithms.

Posterior Density. For small parameter dimensions M and fixed or estimated σ^2 , the posterior density can be constructed by employing quadrature rules to approximate the denominator of (6.6).

If we let w^i and ζ^i denote the quadrature weights and points and assume a noninformative prior, the posterior can be approximated by

$$\pi(\theta|d) \approx \frac{1}{\sum_{i=1}^{N_q} e^{-(SS_{\zeta^i} - SS_q)/2\sigma^2} w^i}. \quad (6.8)$$

The algebraic reformulation of the sum of squares in the denominator is made to avoid numerical 0/0 errors. For $M = 1$ through roughly 4, one can employ tensored Gaussian quadrature relations whereas sparse grid techniques can be employed for moderate dimensionality. We refer to this as *direct Bayesian calibration*.

6.2.2 Sampling Based Metropolis Algorithms

The difficulties associated with approximating the denominator of (6.6) or constructing marginal posterior densities, for moderate to large dimensional parameter spaces, can be partially circumvented by employing Markov chain Monte Carlo (MCMC) techniques. The goal with these algorithms is to construct sampling-based chains whose stationary distribution is the posterior distribution. The capabilities provided in Dakota and QUESO (Quantification of Uncertainty for Estimation, Simulation, and Optimization) are: Delayed Rejection Adaptive Metropolis (DRAM), Differential Evolution Adaptive Metropolis (DREAM), and Gaussian Process Models for Simulation Analysis (GPMSA).

Delayed Rejection Adaptive Metropolis (DRAM)

In Metropolis algorithms, parameters are sampled using a proposal function that reflects, to the degree possible, the geometry of the unknown posterior density. For example, one can propose candidate samples $\theta^* \sim \mathcal{N}(\theta^{k-1}, V)$, where θ^{k-1} is the previous chain element and V is the proposal covariance matrix. Proposed candidates are rejected or accepted with a probability that reflects the degree to which candidates increase the likelihood. The goal is to construct a chain whose stationary distribution is the posterior density.

DRAM is a variation of the Metropolis algorithm whose robustness is improved in two ways. First, adaptation allows the algorithm to update the proposal covariance matrix to reflect accepted candidates. In this manner, information acquired about the posterior distribution through accepted chain candidates is used to update the proposal distribution. Secondly, delayed rejection provides a mechanism for efficiently constructing alternative candidates when the current candidate is rejected. In combination, these two mechanisms provide the algorithm with substantial robustness and efficiency [15]. We note that parallel versions of DRAM have recently been developed [35].

Differential Evolution Adaptive Metropolis (DREAM)

There are various regimes for which DRAM algorithms are often not efficient. These include problems in which posterior densities are multi-modal, highly complex, or have heavy tails. For these cases, the single DRAM chain will be slow to traverse the posterior which can significantly diminish its efficiency. Moreover, the computational overhead associated with complex models can preclude the construction of burned-in single chains whereas one can often compute shorter parallel chains using massively parallel architectures.

DREAM algorithms can circumvent some of these limitations. In these algorithms, candidates are randomly generated using differential evolution algorithms. These algorithms are inherently

parallel and have the advantage that chains can learn from each other. Details are provided in [40, 39].

Gaussian Process Models for Simulation Analysis (GPMSA)

GPMSA provides additional capability for Bayesian calibration. A key part of GPMSA is the construction of a surrogate model or emulator from simulation runs collected at various settings of input parameters; see Chapter 4 for details regarding the construction of surrogate models including Gaussian process representations. The emulator is a statistical model of the system response, and it is used to incorporate the observational data to improve system predictions and constrain or calibrate the unknown parameters. The GPMSA code draws heavily on the theory developed in the seminal Bayesian calibration paper by Kennedy and O’Hagan [22]. The particular approach implemented in QUESO is described in [19]. GPMSA uses Gaussian process models in the emulation, but for functional responses the emulator is actually a set of basis functions (e.g., from a singular value decomposition) which have GPs as the coefficients. One major difference between GPMSA and the QUESO implementation in Dakota is that the QUESO implementation does not have an explicit “discrepancy” function δ which models the difference between the simulation and the observational data results in addition to the error term ϵ , but GPMSA has a sophisticated model for the discrepancy term. The GPMSA implementation in Dakota is presently an early prototype. Details illustrating the performance of GPMSA for COBRA-TF are provided in Section 7.3.

Table 6.4 provides a reference for choosing a Bayesian model calibration method based on the properties of the model.

Table 6.4: Guidelines for Bayesian method selection.

Method Classification	Desired Problem Characteristics	Applicable Methods
DRAM	Unimodal or weakly multi-modal posterior densities; computationally efficient models or surrogates	bayes_calibration queso
DREAM	Multi-modal, complex or heavy tailed posterior densities; inherently parallel simulation codes	bayes_calibration dream
GPMSA	Bayesian calibration using a Gaussian process emulator; can accommodate certain model discrepancy relations	

6.2.3 Model Calibration and Surrogate Models

Dakota provides various capabilities for combining Bayesian model calibration algorithms with surrogate models. The first is to employ a surrogate model or emulator when constructing chains and posterior densities for inputs. For computationally intense codes, this will be necessary to acquire the 10^3 to 10^5 model solutions required to burn-in chains and obtain a statistically relevant number of chain elements. The QUESO GPMSA package provides one alternative for Bayesian calibration using a Gaussian process emulator. Alternatively, the techniques of Chapter 4 can be used to construct a surrogate model, which is then employed in Dakota-QUESO DRAM or Dakota DREAM

for Bayesian model calibration. One can apply Bayesian methods to a general surrogate by using a `model_pointer` to point the Bayesian routine to a model of type `surrogate`. For the nonintrusive PCE or stochastic collocation surrogates, discussed in Section 6.1, or Gaussian process surrogates, one can shortcut this process by specifying the emulator as one of `gaussian_process`, `pce`, etc.

Once input distributions have been constructed, one can employ a surrogate model to implement the sampling methods discussed in Section 6.1.1 for quantities of interest that may not have been used for calibration. Alternatively, one can employ the stochastic polynomial methods of Section 6.1.2 to propagate uncertainties. We note that these comprise a form of interpolation or regression-based surrogate models, which complement the kriging and Gaussian process-based methods detailed in Chapter 4.

6.2.4 Verification

We summarize here a general framework for verifying model calibration results for CASL codes implemented in VERA via Dakota. This framework is generally applicable to codes with nonlinear parameter dependencies and experimental or synthetic data.

- (i) Test algorithms using a linearly-parameterized model where analytic uncertainty relations can be computed. Whereas this is not generally possible for CASL codes with nonlinear parameter dependencies, code verification in this manner provides a first step for verifying the capabilities of the model calibration framework and it may be used in certain nearly linear operating regimes.
- (ii) Compare to direct numerical implementation of Bayes' formula (6.6) for small to moderate input or parameter dimensions M ; e.g., $M \leq 20$ to 30. For the likelihood relation (6.7) and a noninformative prior $\pi_0(\theta)$, this involves the evaluation of the relation (6.8). This also comprises code verification.
- (iii) Compare to other packages that implement DRAM to perform code-to-code verification.
- (iv) Compare DRAM and DREAM results.
- (v) Compare to sampling distributions provided by frequentist analysis. Whereas this approach can guide verification, it must be used with care since the underlying assumptions for frequentist and Bayesian inference differ significantly; see Chapter 4 of [34]. For example, asymptotic analysis often yields Gaussian sampling distributions which will obviously be inaccurate if the true distribution is highly non-Gaussian.
- (vi) Check the convergence of the algorithms by increasing the number of quadrature points used in (6.8) or number of iterations in DRAM or DREAM chains to establish solution verification.

6.2.5 Synthetic Data

One often employs synthetic data when testing model calibration algorithms since it provides a regime in which errors are constructed, and hence known, and it can be employed when experimental data is not readily available. We illustrate the construction of synthetic data for the statistical model (6.5) but note that the procedure may vary for other statistical models.

For a nominal input $\tilde{\theta} = [\tilde{\theta}_1, \dots, \tilde{\theta}_M]$, one calculates a nominal model response $y(\tilde{\theta})$. For a specified variance σ^2 , one then generates realizations $\epsilon_1, \dots, \epsilon_N$ from a normal distribution $\mathcal{N}(0, \sigma^2)$, which yields the synthetic data

$$d_i = y(\tilde{\theta}) + \epsilon_i, \quad i = 1, \dots, N.$$

For multiple responses, the standard deviation σ is typically scaled by the magnitude of each response in the manner illustrated in the cantilever beam example of Section 6.2.6.

6.2.6 Bayesian Calibration Examples

We illustrate here the performance and verification of the Dakota Bayesian model calibration packages for the cantilever beam example of Section 6.2.4 and linear verification example of Section 2.2 and Appendix A. The use of the algorithms for COBRA-TF is illustrated in Section 7.3.

Cantilever Beam

We employ the cantilever beam example of Section 2.1 to illustrate the implementation of the DRAM and DREAM algorithms for Bayesian model calibration along with the verification framework summarized in Section 6.2.4. The posterior parameter densities constructed in this manner can subsequently be employed with the uncertainty propagation techniques detailed in Section 6.1 to quantify response uncertainties.

Case I. We consider first the case in which the Young's modulus E and width w are considered unknown, the remaining parameters and inputs t, R, L, D_0, X and Y are assumed known and fixed, and data are taken to be displacement and stress measurements d and s . To construct synthetic data using the techniques of Section 6.2.5, based on the assumption of independent and identically distributed (iid) observation errors $\epsilon_i \sim \mathcal{N}(0, \sigma^2)$, we compute the nominal displacement and stress values

$$\tilde{d} = 0.3086, \tilde{s} = 2.6667 \times 10^3$$

using the fixed and nominal input values in Table 6.5. The standard deviations are taken to be

$$\sigma_d = 0.1 \cdot \tilde{d}, \quad \sigma_s = 0.1 \cdot \tilde{s}$$

which yields the synthetic data compiled in Table 6.6 when observation errors ϵ_{d_i} and ϵ_{s_i} are drawn from normal distributions $\mathcal{N}(0, \sigma_d^2)$ and $\mathcal{N}(0, \sigma_s^2)$.

t	R	L	D_0	X	Y	$\tilde{E}$	$\tilde{w}$
3	4×10^4	100	2.2535	500	100	2.85×10^7	2.5

Table 6.5: Known values for t, R, L, D_0, X, Y and nominal values $\tilde{E}$ and $\tilde{w}$.

$d (\times 10^{-1})$	3.2075	2.7005	2.7939	2.8578	2.9298	2.9875	3.0903	2.1515	2.9454	3.4700
$s (\times 10^3)$	2.3838	3.0943	2.9959	2.6054	2.2650	2.5481	2.6251	2.7403	2.5970	2.7849

Table 6.6: Synthetic displacement and stress data employed for Bayesian model calibration.

We first employ the discretized Bayes relation (6.8) to directly construct marginal posterior densities for E and w . Due to the simplicity of the algebraic model, we employ a tensored trapezoid quadrature rule, which yields the convergence results compiled in Table 6.7. For more computationally intensive models and codes, one would employ tensored Gaussian routines for low parameter dimensions for $M = 1$ to approximately 6 and sparse grid techniques for moderate dimensionality of M up to 30 or 40 where the upper limit depends on the regularity of the likelihood.

N_q	μ_E	σ_E	μ_w	σ_w
40	2.8751e+07	2.6685e+05	2.5024	1.2589e-02
80	2.8688e+07	1.6430e+05	2.5004	6.3164e-03
160	2.8693e+07	1.3087e+05	2.5001	3.6669e-03
320	2.8693e+07	1.3011e+05	2.5001	3.5940e-03

Table 6.7: Convergence of the direct numerical Bayes relation.

The direct results are compared with posterior densities constructed using the Dakota-QUESO DRAM and Dakota DREAM algorithms in Figure 6.5. The corresponding input decks are provided in Listings 6.4 and 6.5 and the means and standard deviations are compiled in Table 6.8. The joint sample points plotted in Figure 6.6 demonstrate that E and w are correlated but identifiable. For the DRAM algorithm, we constructed a chain of length 50,000 to ensure burn-in, whereas we employed 10 chains of length 5,000 in the DREAM implementation. In both cases, we employed noninformative priors. The matching of posterior densities and moments with those constructed through direct computation verifies the accuracy of the DRAM and DREAM implementations. *For CASL applications, we recommend that, when possible, at least two of the techniques be compared to verify the accuracy of the inference procedure.*

	μ_E	σ_E	μ_w	σ_w
Direct	2.8693e+07	1.3011e+05	2.5001	3.5940e-03
DRAM	2.8694e+07	1.3101e+05	2.5002	3.5894e-03
DREAM	2.8697e+07	1.2866e+05	2.5002	3.5540e-03

Table 6.8: Posterior means and standard deviations provided by the direct, DRAM and DREAM algorithms.

Case II. Secondly, we consider the case when E , t and w are considered uncertain and synthetic data is taken to be displacement, stress and area measurements generated in a manner analogous to Case I. The joint sample points constructed using Dakota-QUESO DRAM are plotted in Figure 6.7. The nearly single-valued relation between t and w indicates that these parameters are essentially nonidentifiable, which is a manifestation of nonunique input-output maps. This is consistent with the observation that the product $A = w \cdot t$ appears in the displacement, stress and area relations. The implementation of Bayesian calibration techniques for nonidentifiable parameter sets will generally be problematic unless informative prior specification is provided. If such prior information is not available, parameter selection based on global sensitivity analysis or reduced order modeling

Listing 6.4: Input for Dakota-QUESO DRAM.

```

1 # DAKOTA INPUT FILE - cantilever_bayes.in

3 strategy,
    single_method
5     tabular_graphics_data

7 method,
    bayes_calibration queso,
9     samples = 50000 seed = 348
    mcmc_type dram
11    rejection delayed
    metropolis adaptive
13    output verbose

15 variables,
    proposal_covariance_scale 1.e10 1.e3
17    active design
    continuous_design 2
19    upper_bounds 1.e8 10.0
    initial_point 2.85e7 2.5
21    lower_bounds 1.e6 0.1
    descriptors 'E' 'w'
23    continuous_state 4
    initial_state 3 40000 500 1000
25    descriptors 't' 'R' 'X' 'Y'

27 interface,
    system
29    analysis_driver = 'cantilever2'

31 responses,
    calibration_terms = 2
33    calibration_data_file = 'dakota_cant2_withsigma.dat'
    freeform
35    num_experiments = 1
    num_replicates = 10
37    num_std_deviations = 2
    descriptors = 'stress' 'displacement'
39    no_gradients
    no_hessians

```

Listing 6.5: Input for Dakota DREAM.

```

1 # DAKOTA INPUT FILE - cantilever_bayes.in
3 strategy,
4     single_method
5     tabular_graphics_data
7 method,
8     bayes_calibration dream,
9     samples = 50000 seed = 348
10    chains = 10
11    output verbose
13 variables,
14     active design
15     continuous_design 2
16     upper_bounds 1.e8 10.0
17 # Truth value is 2.85e7 2.5
18     lower_bounds 1.e6 0.1
19     descriptors 'E' 'w'
20     continuous_state 4
21     initial_state 3 40000 500 1000
22     descriptors 't' 'R' 'X' 'Y'
23
24 interface,
25     system
26     analysis_driver = 'cantilever2'
27
28 responses,
29     calibration_terms = 2
30     calibration_data_file = 'dakota_cant2_withsigma.dat'
31     freeform
32     num_experiments = 1
33     num_replicates = 10
34     num_std_deviations = 2
35     descriptors = 'stress' 'displacement'
36     no_gradients
37     no_hessians

```

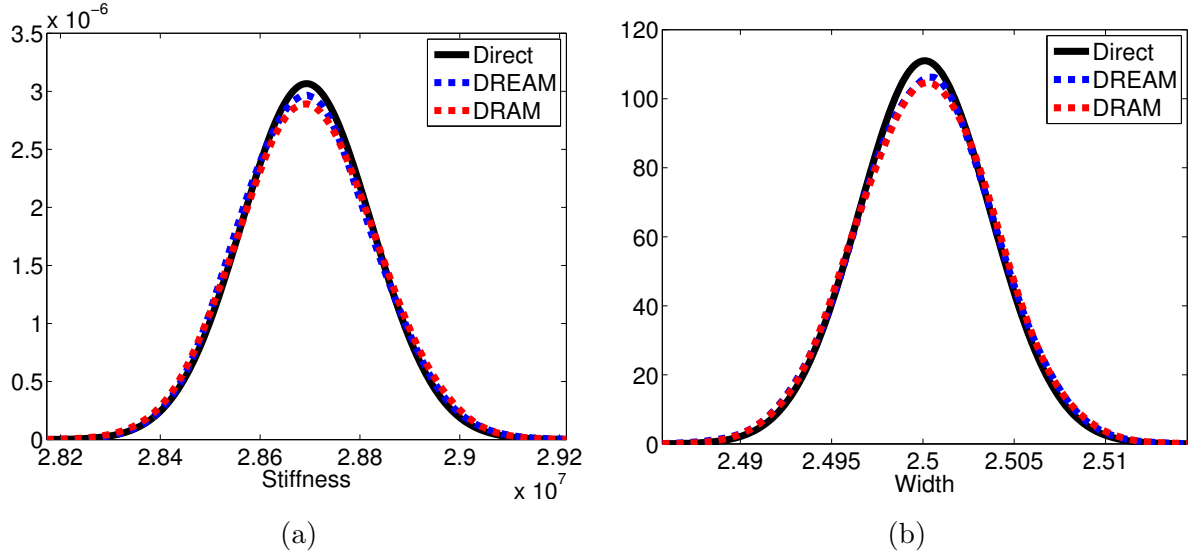



Figure 6.5: Marginal posterior densities for (a) E and (b) w generated through direct solution of Bayes' relation (6.8), Dakota-QUESO DRAM and Dakota DREAM.

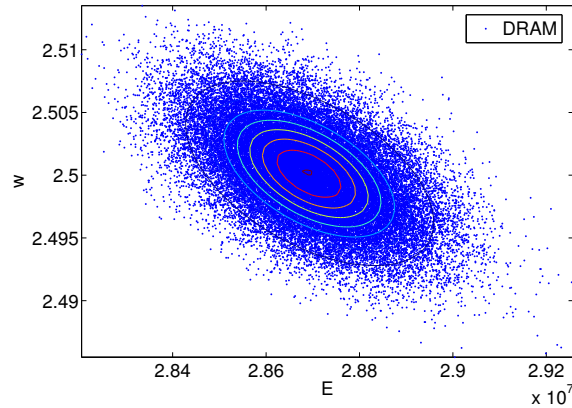


Figure 6.6: Joint posterior sample points for E and w constructed using the direct method (contours) and DRAM (points).

techniques should be employed to determine the set of identifiable or influential parameters. Details regarding parameter selection techniques can be found in Chapter 6 of [34].

General Linear Model Verification Tests

We illustrate here aspects of the general linear model verification test suite detailed in Chapter 2 and Appendix A. This constitutes step (i) in the verification framework detailed in Section 6.2.4. All results were obtained using the QUESO implementation of DRAM.

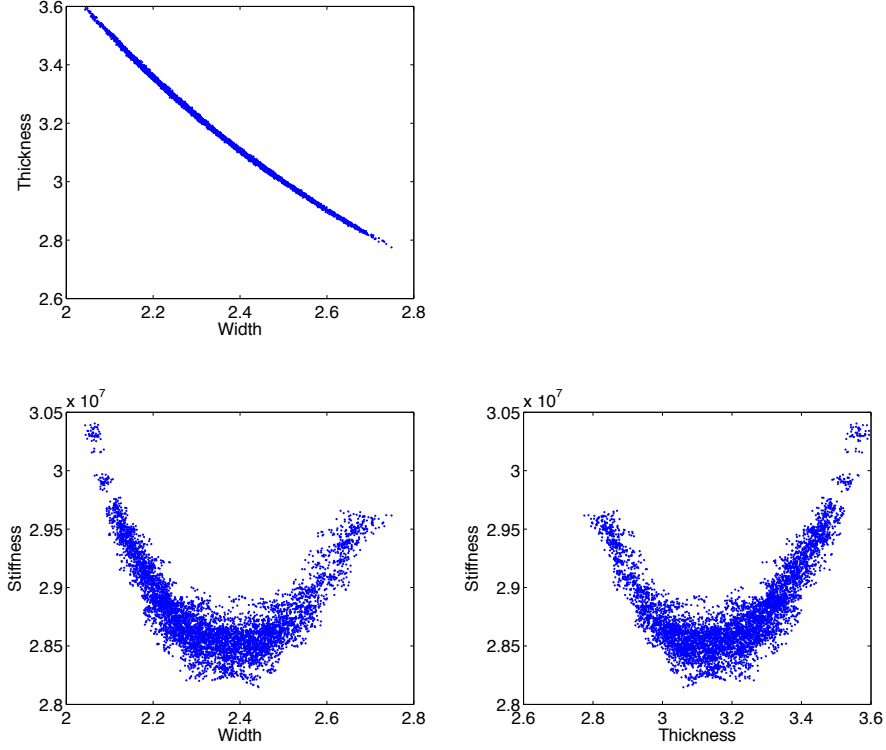


Figure 6.7: Joint posterior sample points for E, w and t using DRAM.

The true values for parameters were constructed using the relations

$$\beta_0^i = 0.4 + 0.05 \sin \left(2\pi \frac{i-1}{N_\beta} \right), \quad i = 1, \dots, N_\beta,$$

$$\lambda_0 = 1000, \quad \phi_0 = 0.3.$$

The superscript i designates the components of β_0 .

Example 1. We first generate observations via (2.4) assuming uncorrelated noise, and then calibrate the unknown variables for several choices of sample size and two separate cases of unknown parameters. This allows verification of convergence of the calibrated variables β and (where applicable) λ to their true values. We consider the cases:

- Case 1 (β unknown), $N_\beta = 1$, $N \in \{1, 2, 3, 10, 10^2, 10^3, 10^4, 10^5, 10^6\}$,
- Case 2 (β, λ unknown), $N_\beta = 1$, $N \in \{2, 3, 10, 10^2, 10^3, 10^4, 10^5, 10^6\}$,
- Case 1 (β unknown), $N_\beta = 10$, $N \in \{10, 10^2, 10^3, 10^4\}$,
- Case 2 (β, λ unknown), $N_\beta = 10$, $N \in \{15, 10^2, 10^3, 10^4\}$.

Recall that specifying β as unknown is equivalent to stating that β is calibrated to the data. The $N_\beta = 1$ cases were computed using a DRAM chain with 20,000 iterates and the $N_\beta = 10$ cases used 50,000 iterates.

To facilitate visualization for the $N_\beta = 10$ case and provide an intuitive sense of scale, the formulas

$$\begin{aligned}\tilde{\beta}_m &= 100 \frac{\|E(\beta) - \beta_0\|_2}{\|\beta_0\|_2}, & \tilde{\beta}_v &= 100 \frac{\sqrt{\text{var}(\beta)}}{\|\beta_0\|_2}, \\ \tilde{\lambda}_m &= 100 \frac{\|E(\lambda) - \lambda_0\|_2}{\|\lambda_0\|_2}, & \tilde{\lambda}_v &= 100 \frac{\sqrt{\text{var}(\lambda)}}{\|\lambda_0\|_2}\end{aligned}$$

were used to normalize the results in the plots. Here $E(\cdot)$ denotes the expectation or mean of the chain, $\text{var}(\cdot)$ denotes the variance of the chain, and $\|\cdot\|_2$ is the Euclidean norm of the argument. The variables with the m subscript are errors in the expectations of the calibrated variables relative to true values represented by the variables with subscript 0. The variables with the subscript v are standard deviations of the chains relative to the magnitude of the true parameters, which provides a more independent sense of the estimated uncertainty. Both of these variables are expressed as percentages and are expected to converge to 0 with increasing sample size for uncorrelated observation error. Figure 6.8 illustrates the results of the convergence studies plotted on logarithmic axes. The downward trend of the plots confirms convergence of the computed quantities to 0 for all considered cases.

We re-emphasize that the results obtained using DRAM are a sequence of iterates drawn from a distribution that characterizes the confidence in the estimate of the unknown parameters. To illustrate this point, histograms of the iterates obtained for Case 1 (unknown β) assuming a non-informative prior are shown in Figure 6.9 for various sample sizes N . For this problem, the β posterior densities are expected to be Gaussian as confirmed by the shapes of the histograms. As the sample size N increases, the mean of the Gaussian posterior density approaches the true parameter $\beta_0 = 0.4$ while the variance of the posterior density decreases. A Gaussian density converges to a Dirac delta distribution (sometimes called an impulse function) as its variance decreases which Figure 6.9 illustrates. This allows us to conclude that for the case where β is unknown, assigned a noninformative prior, and calibrated to observations corrupted by uncorrelated noise, the estimate for β tends towards the deterministic value β_0 with increasing sample size N . This is consistent with the intuition that confidence in an estimate increases as the number of data measurements increases.

Example 2. Here we compare the posterior densities for the calibrated variables obtained via DRAM with the exact expressions for the posterior densities computed analytically in Appendix A. For this example, we used a sample size of $N = 100$ observations generated with equi-correlated errors, and considered $N_\beta = 1$ regression parameter. We computed the following cases:

- Case 1 (β unknown) with noninformative prior,
- Case 1 (β unknown) with Gaussian prior,
- Case 2 (β, λ unknown) with noninformative prior.

For the Gaussian prior case, we set the mean of the prior to $\mu_0 = 0.6$ and the variance to 2×10^{-4} . The mean was chosen relatively far from the true value of $\beta_0 = 0.4$ to distinguish the results from those of the noninformative case, thus reducing the possibility of coincidentally correct output. All cases were computed with parameter chains of 20,000 iterates.

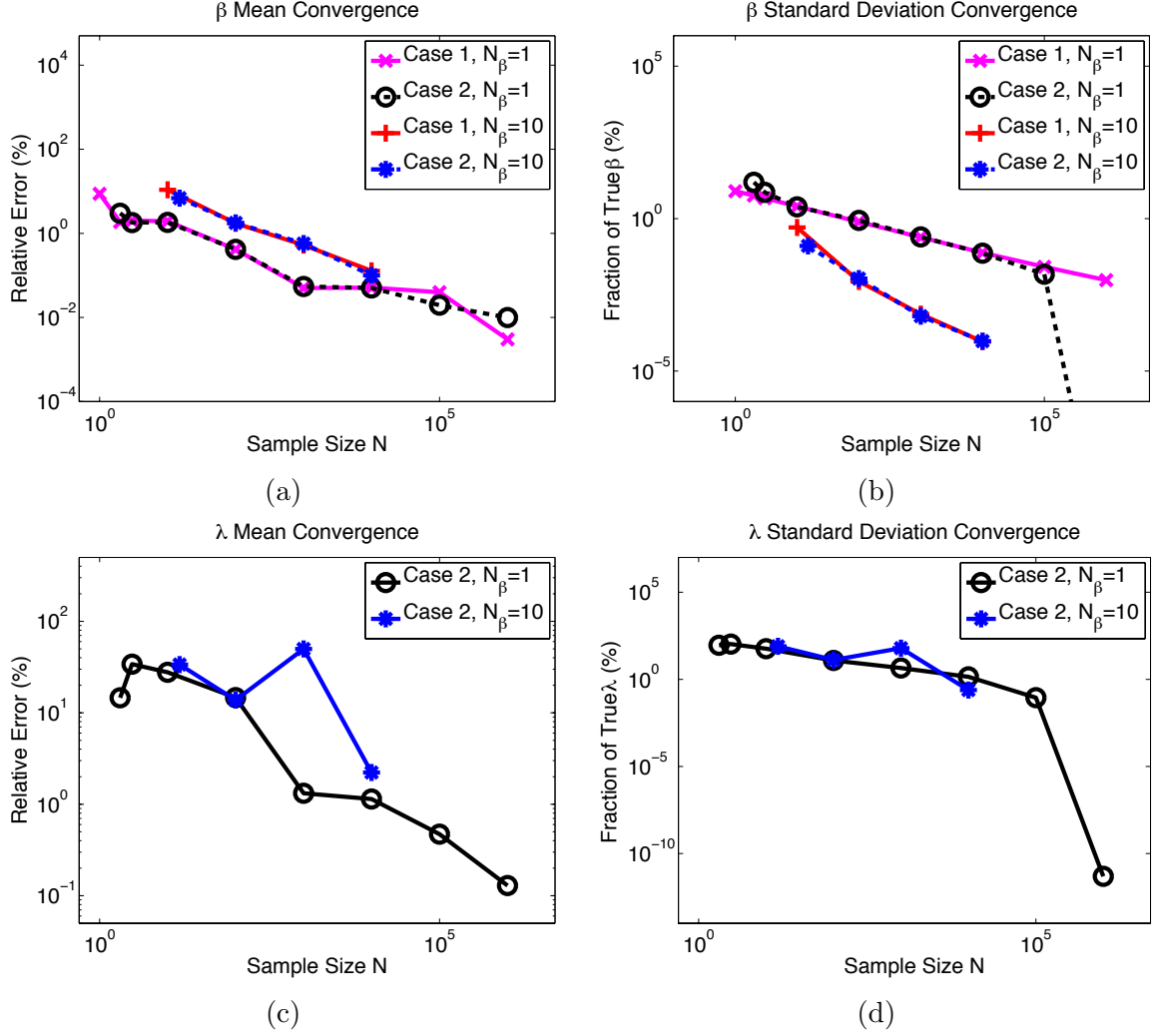


Figure 6.8: Convergence studies for the general linear model verification test suite with uncorrelated noise and noninformative priors. (a) Percentage error in the mean of β relative to β_0 . (b) Standard deviation of β expressed as percentage of β_0 . (c) Percentage error in mean of λ relative to λ_0 , and (d) Standard deviation of λ expressed as percentage of λ_0 .

In Figure 6.10, we show the results for each case. As detailed in Appendix A, Case 1 yields a Gaussian posterior density for β whereas the posterior for β in Case 2 is a t -distribution. However, since $N = 100$ is a moderate sample size (i.e., not small), the t -distribution is roughly Gaussian. The posterior for λ is a Gamma distribution. The DRAM-computed results show good agreement with the analytic results.

Example 3. Here we compare the analytically computed posterior densities with those obtained from DRAM for small sample sizes (i.e., $N \in \{2, 3, 4, 10\}$) and an unknown λ to verify that the posterior of β matches the expected t -distribution. Both β and λ were calibrated and the parameter chains were computed to 20,000 iterates. We computed the calibration for the three cases of uncorrelated observation errors, equi-correlated observation errors, and AR(1) correlated

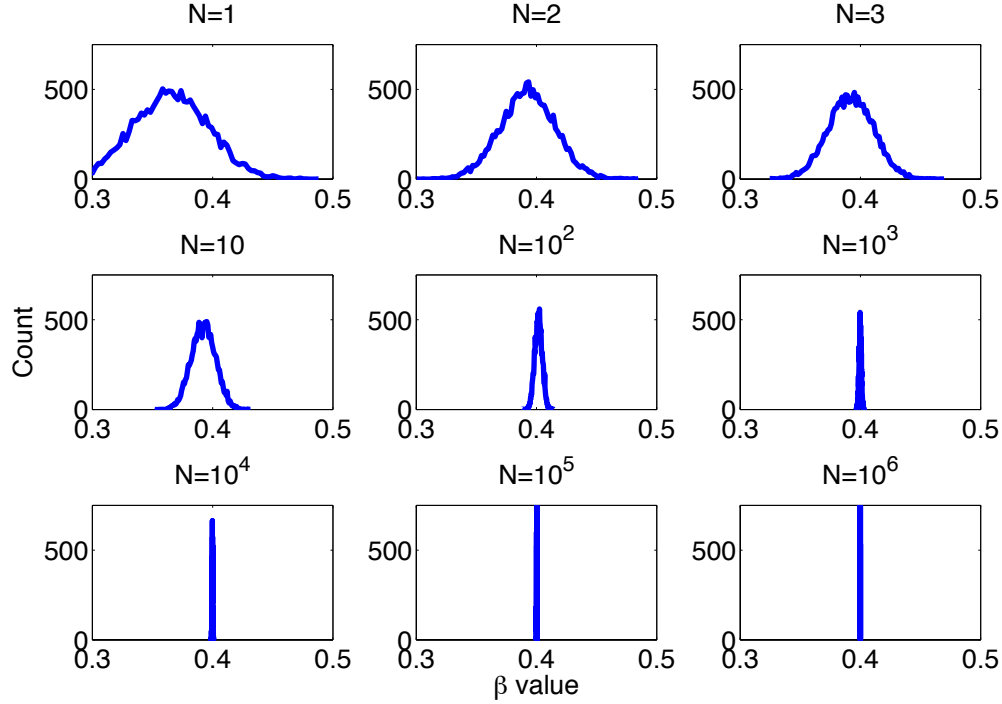


Figure 6.9: Histograms of the chains obtained from DRAM with β calibrated assuming a noninformative prior. The figures illustrate convergence of the posterior densities to a Dirac delta impulse centered at the true value $\beta_0 = 0.4$ as the sample size N increases.

observation errors as specified by the ridge functions detailed in Appendix A.2.

The results are shown in Figure 6.11 for uncorrelated observation errors, Figure 6.12 for equi-correlated observation errors, and Figure 6.13 for AR(1) observation errors. In each of the plots for the β posterior densities, we also plot the Gaussian (labeled “Normal” in the figures) posterior density which would be obtained for known λ to illustrate convergence of the t -distribution to a Gaussian as the sample size N increases. The DRAM-computed results again generally agree with the analytic results.

A sample QUESO input deck is provided in Listing 6.6. This general linear model verification test suite will be integrated into future Dakota releases.

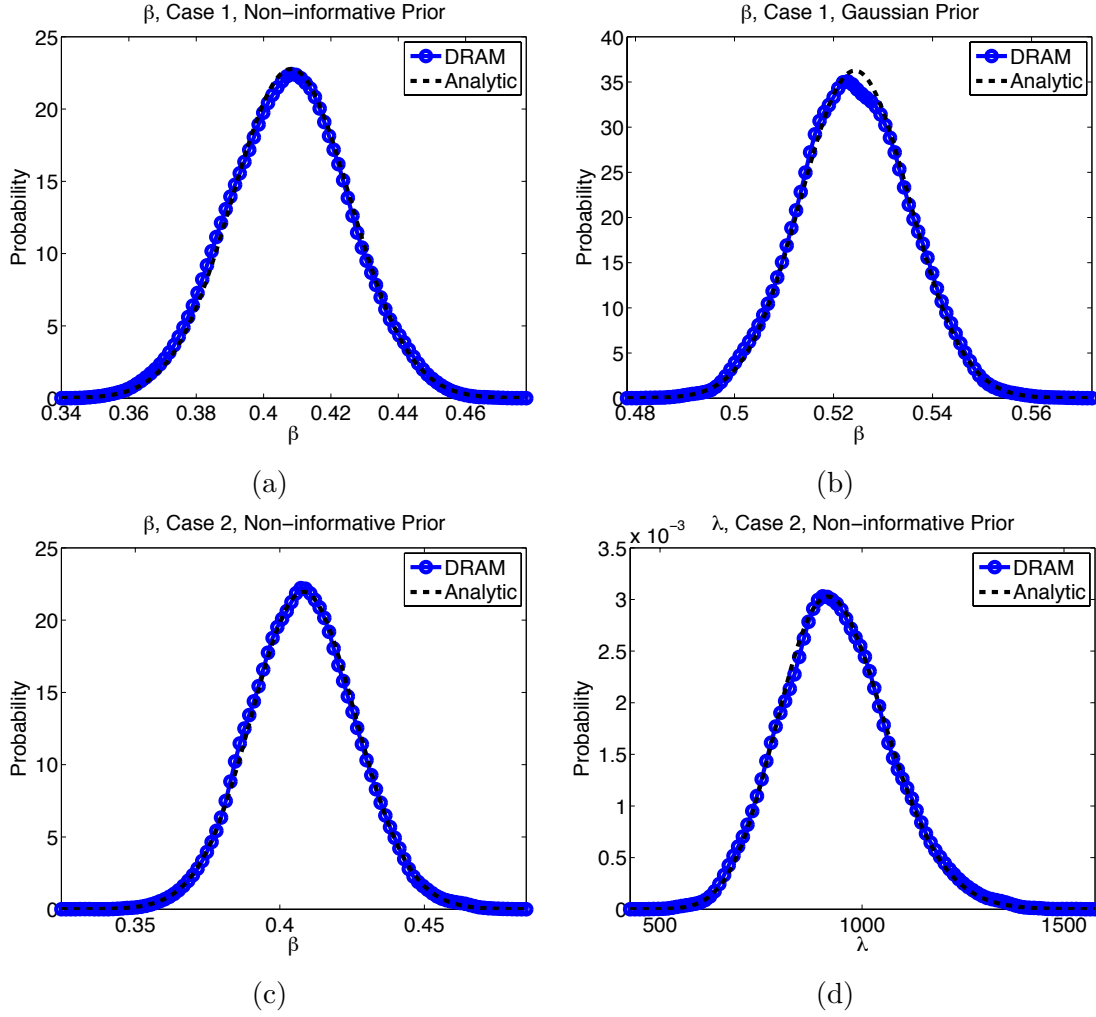


Figure 6.10: Comparison of DRAM-computed posterior densities for β, λ with the analytic posterior densities. Case 1: β posterior density with β unknown assuming a (a) noninformative prior and (b) Gaussian prior. Case 2: β, λ unknown, calibrated assuming a noninformative prior, and the resulting (c) β posterior density and (d) λ posterior density.

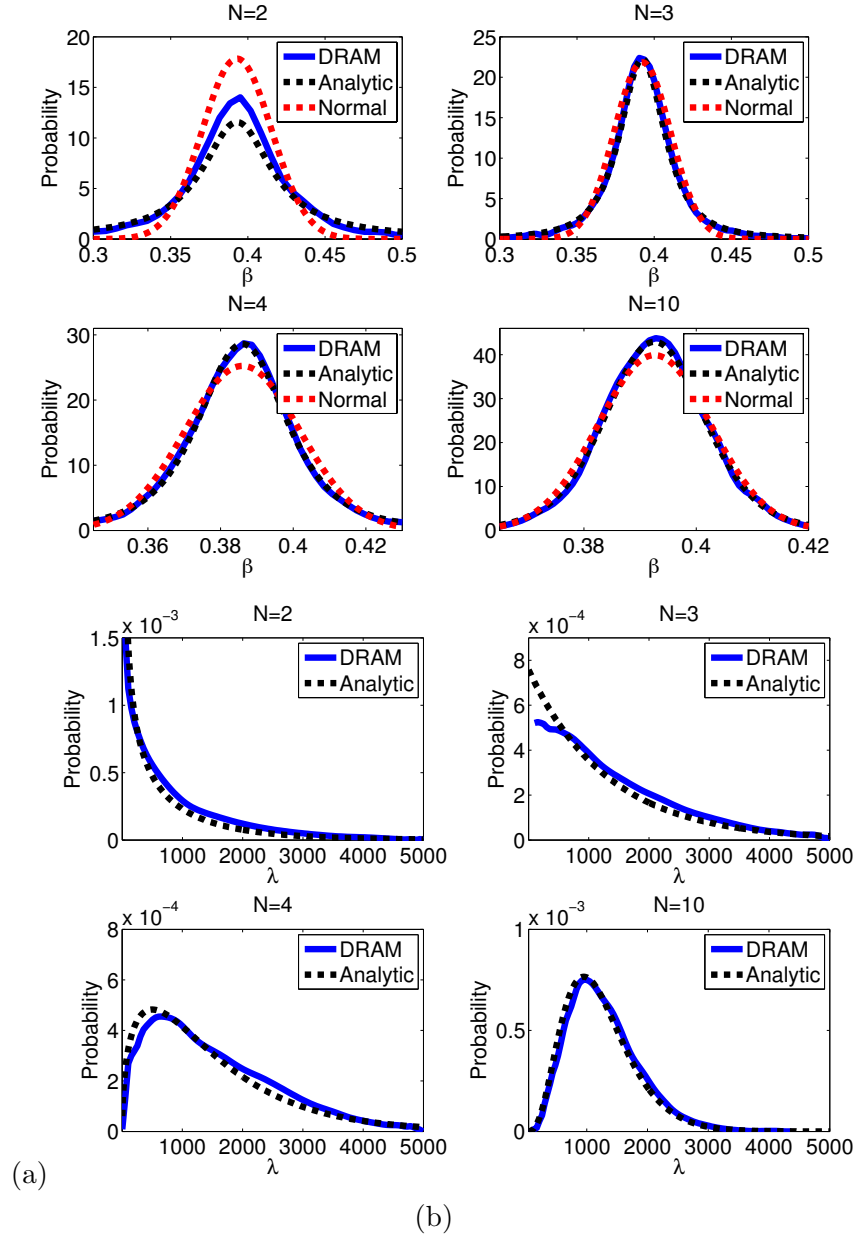
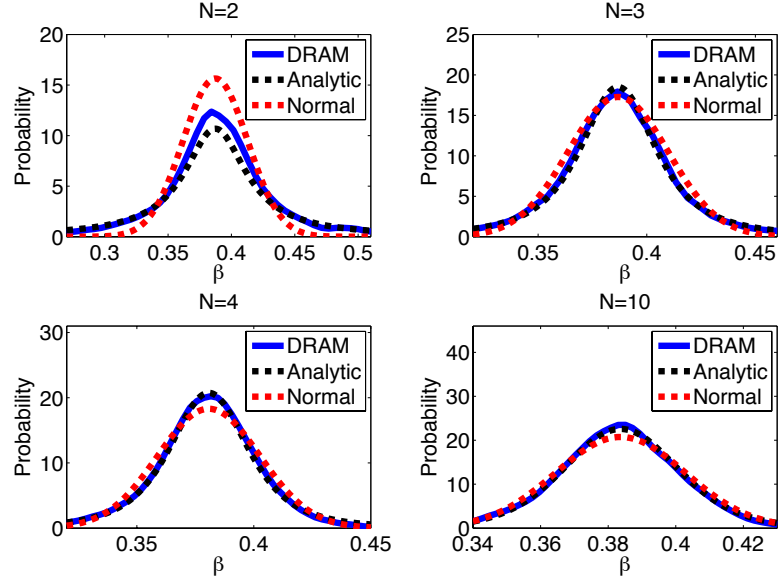
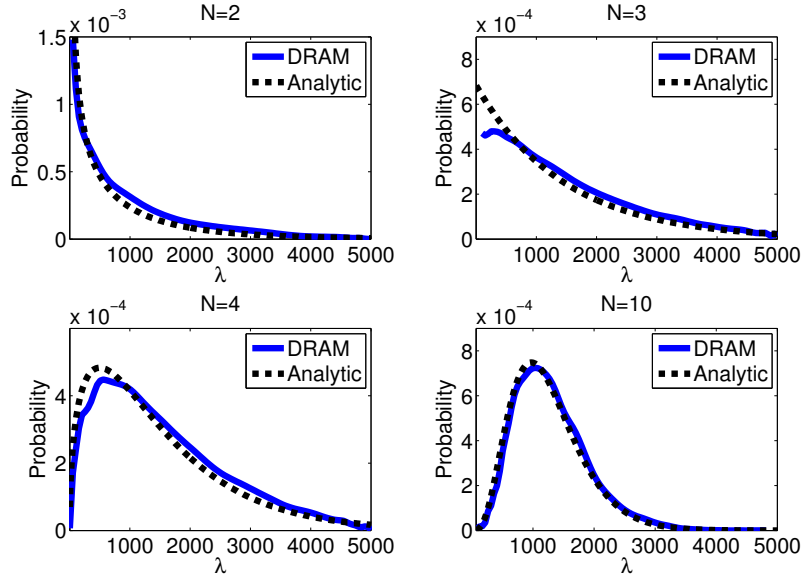


Figure 6.11: Comparison of DRAM-computed posterior densities with analytic posterior densities for small sample sizes and uncorrelated observation errors. Normal distribution shown to illustrate convergence of the t -distribution with increasing sample size. (a) Posterior densities for β and (b) Posterior densities for λ .

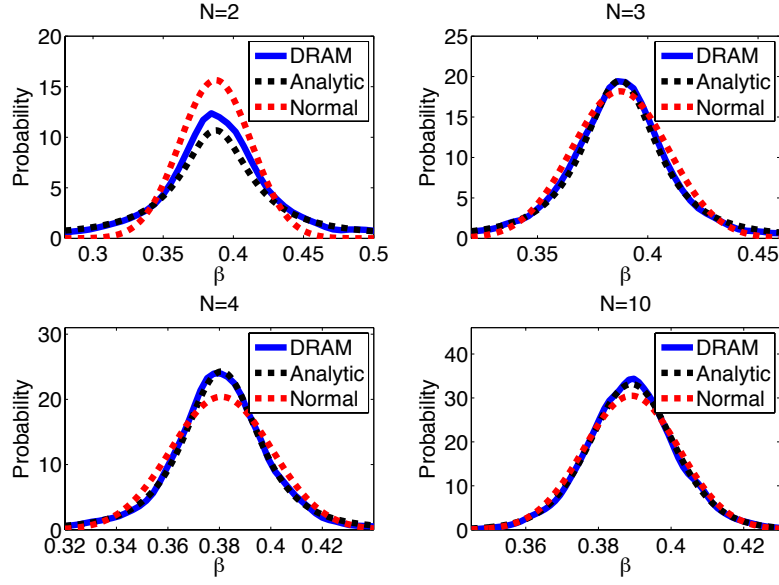


(a)

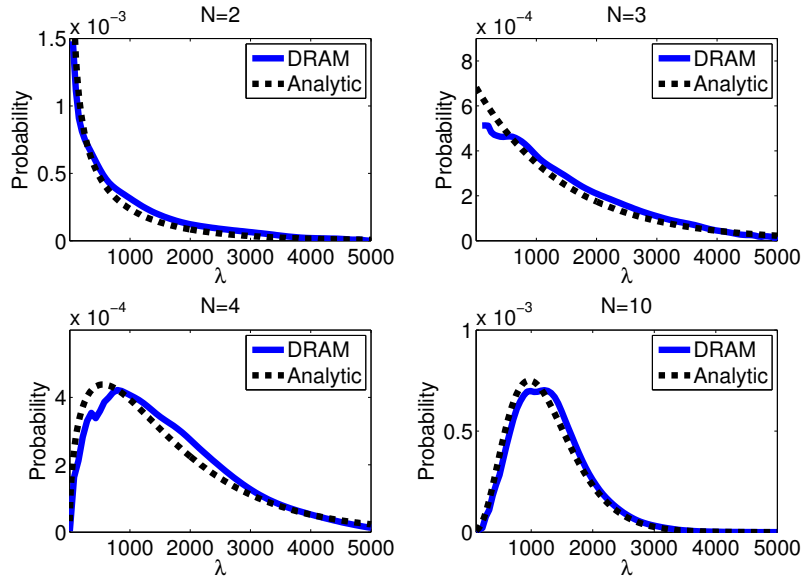


(b)

Figure 6.12: Comparison of DRAM-computed posterior densities with analytic posterior densities for small sample sizes and equi-correlated observation errors. Normal distribution shown to illustrate convergence of the t -distribution with increasing sample size. (a) Posterior densities for β and (b) Posterior densities for λ .



(a)



(b)

Figure 6.13: Comparison of DRAM-computed posterior densities with analytic posterior densities for small sample sizes and AR(1) observation errors. Normal distribution shown to illustrate convergence of the t -distribution with increasing sample size. (a) Posterior densities for β and (b) Posterior densities for λ .

Listing 6.6: QUESO input for the general linear model verification test suite.

```

#####
2 # UQ Environment
#####
4 env_numSubEnvironments      = 1
  env_subDisplayFileName     = outputData/display_env
6 env_subDisplayAllowAll     = 0
  env_subDisplayAllowedSet   = 0 1 2 3 4 5 6 7
8 env_displayVerbosity       = 2
  env_seed                   = 0
10
#####
12 # Statistical inverse problem (ip)
#####
14 ip_computeSolution         = 1
  ip_dataOutputFileName      = outputData/sip_algver
16 ip_dataOutputAllowedSet   = 0 1
#####
18 # Information for Metropolis-Hastings algorithm
20 #####
  ip_mh_dataOutputFileName   = outputData/sip_algver
22 ip_mh_dataOutputAllowedSet = 0 1
#####
24 ip_mh_rawChain_dataInputFileName = .
  ip_mh_rawChain_size        = 20000
26 ip_mh_rawChain_generateExtra    = 0
  ip_mh_rawChain_displayPeriod    = 2000
28 ip_mh_rawChain_measureRunTimes  = 1
  ip_mh_rawChain_dataOutputFileName = outputData/sip_algver_raw_chain
30 ip_mh_rawChain_dataOutputAllowedSet = 0 1 2 3 4 5 6 7
#####
32 ip_mh_displayCandidates          = 0
  ip_mh_putOutOfBoundsInChain      = 0
34 ip_mh_dr_maxNumExtraStages       = 3
  ip_mh_dr_listOfScalesForExtraStages = 5. 10. 20.
36 ip_mh_am_initialNonAdaptInterval = 0
  ip_mh_am_adaptInterval           = 200
38 ip_mh_am_eta                     = 0.05703    #(2.4^2)/d, d is
                                           the dimension of the problem
40 ip_mh_am_epsilon                 = 1.e-5
#####
42 ip_mh_filteredChain_generate      = 1
  ip_mh_filteredChain_discardedPortion = 0.
44 ip_mh_filteredChain_lag          = 10
  ip_mh_filteredChain_dataOutputFileName = outputData/sip_algver_filtered_chain
46 ip_mh_filteredChain_dataOutputAllowedSet = 0 1

```

Chapter 7

COBRA-TF VUQ Studies

This chapter concludes the manual by demonstrating use of Dakota to complete an overall VUQ process for CASL Progression Problem 6, simulated with COBRA-TF, as described in Section 2.3. The workflow demonstration includes the following Dakota studies:

1. Initial **centered parameter studies** to exercise the COBRA-TF model with two coupled physics scenarios and verify the Dakota/COBRA-TF interfaces, resulting in adding and removing some parameters from the admissible set. (Section 7.1)
2. **Sensitivity analysis** using parameter study, LHS and Morris methods to identify the most important of 33 parameters. Initial LHS studies revealed code robustness issues under joint variation, resulting in adjusting the range of one parameter. Screening based on these studies resulted in five significant parameters. Another LHS study with 50 samples was conducted over these parameters and used in subsequent activities. (Section 7.2)
3. **Deterministic and Bayesian calibration** to estimate the values or distributions of the five significant parameters using synthetic data. The deterministic calibration demonstrates gradient-based local calibration using the COBRA-TF simulation model directly, while the Bayesian calibration uses a surrogate model constructed from the 50 LHS samples. (Section 7.3)
4. **Surrogate construction** and validation to assess the quality and applicability of a response surface model. (Section 7.3.2)

The total Progression Problem 6 parameter set for consideration is indicated in Table 2.6 in Section 2.3. The two parameters marked with an asterisk in this table, `q1` and `qv`, were initially included to assess the effects of heat transfer across the fuel pin surface into the channel liquid and vapor phases, respectively. However, it was quickly discovered that for the steady-state Progression Problem 6, perturbing either of these produces an inherent thermal imbalance precluding any steady thermal behavior. Accordingly, these parameters were excluded from the following studies. For all studies in this chapter, the total pressure drop through the fuel rod assembly (here indicated by `TotalPressure`) was used as the quantity of interest.

7.1 Initial Parameter Studies with Two Power Distributions

Two initial centered parameter studies were conducted to verify the Dakota/COBRA-TF interface, assess code robustness, and generate initial results. Both employed a Dakota centered parameter study (Section 3.2.1) over 29 parameters for which the shift values were zero, e.g. $ka_p = 0$, and the scaling values were allowed to vary by $\pm 5\%$ around unity, e.g. $k_p = 1.0 \pm 5\%$ in increments of 1%. The two studies differ in the power distribution input to the thermal hydraulics code. The first uses a uniform value specified via input, while the second uses an axially varying power distribution representing a converged steady-state solution from a previous full simulation of Progression Problem 6. The second study represents a parameter sensitivity study performed around the actual solution to the neutronics component of the problem. The four parameters `k_xkwlx`, `k_cd`, `k_cdfb`, and `k_wkr` were not yet available in the Dakota/COBRA-TF interface when these early studies were conducted.

The initial centered parameter studies each involve 29 parameters evaluated at 10 perturbed values in addition to the baseline (nominal) evaluation. This amounts to 291 total runs for each study. Dakota provides a concurrent execution facility which for these studies enables 60 independent runs to execute simultaneously on the `james007` CASL machine. Each run requires between 5 and 7 minutes so that each centered parameter study completes in just under 1 hour assuming available capacity on the machine.

Table 7.1 summarizes the results for the uniform power distribution, and Table 7.2 summarizes results for the spatially varying power distribution. All values are reported as percentages of the difference in total pressure drop across the fuel assembly compared to a baseline value representing unperturbed parameter values. All parameters were perturbed in the same manner, e.g. the scaling coefficient was adjusted -5% to +5% in increments of 1%. The inlet mass flow rate `finlet` was incorrectly included in these initial studies and later removed in the final study when the axial wall friction and the three loss coefficients for grid spacers were added. These initial results demonstrate the iterative exploratory process of conducting Dakota studies on models, so detailed discussion is omitted.

7.2 COBRA-TF Sensitivity Studies

Initial studies helped stabilize the admissible parameter set to 33 key parameters, for which the VUQ process is demonstrated in this section. The first step in sensitivity analysis is to perform a centered parameter study to assess the effect of individual parameters on the simulation response. Then we conduct Dakota analyses that jointly vary the parameters to better assess global sensitivities for complex models. The Latin hypercube and Morris methods complement the parameter study results to screen the parameter set.

7.2.1 Centered Parameter Study

Dakota Input: A Dakota input file `dakota_centered_33.in` for a 33 variable (now including parameters `k_xkwlx`, `k_cd`, `k_cdfb`, and `k_wkr`) centered parameter study is shown in Listing 7.1. The `method` section (line 6) prescribes the study with five each positive and negative parameter steps of 0.01. The 33 COBRA-TF parameters are specified in the `variables` section, with an `initial_point` of 1.0 (line 14), indicating nominal input values for the simulation. The `interface` (line 17) specifies use of the `dakota-vera-analysis` driver (line 21), which implements the Dakota/COBRA-TF

Table 7.1: Percent difference of total pressure drop compared to the baseline value of 1.17304 bar using uniform power distribution input to the COBRA-TF thermal hydraulics code.

	percent perturbation in parameter									
parameter	-5%	-4%	-3%	-2%	-1%	+1%	+2%	+3%	+4%	+5%
cond	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
eta	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
finlet	-7.07	-5.68	-4.28	-2.87	-1.44	1.45	2.91	4.39	5.88	7.38
gama	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
qliht	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
qradd	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
qradv	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
qvapl	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
rodqq	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
sdent	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
sent	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
sphts	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
tmasg	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
tmasl	-0.002	-0.002	-0.001	-0.001	-0.001	0.0	0.001	0.001	0.002	0.003
tmasv	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
tmome	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
tmoml	0.003	0.002	0.002	0.001	0.0	-0.001	-0.003	-0.003	-0.005	-0.006
tmomv	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
tnrgl	0.0	0.0	0.0	0.0	0.0	0.0	-0.001	-0.001	-0.001	-0.001
tnrgv	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
xk	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
xkes	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
xkge	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
xkl	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
xkle	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
xkvls	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
xkwew	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
xkwlw	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
xkwvw	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0

interface described in Appendix B. This analysis workflow accepts values of the 33 parameters from Dakota, runs the simulation, and extracts the desired response metric, **TotalPressure**. Dakota's **work_directory** feature (line 28) will cloister each COBRA-TF simulation in a separate working directory to permit concurrent model evaluations.

Results and Discussion: Figure 7.1 displays data from the Dakota-generated tabular data file **dakota_centered_33.dat**, revealing the univariate effects of each parameter on the **TotalPressure** response. The parameters **k_cd** and **k_xkw1x** have significant effect and there is a strong linear re-

Table 7.2: Percent difference of total pressure drop compared to the baseline value of 1.17632 bar using a power distribution from a previous steady-state neutronics solution to Progression Problem 6.

	percent perturbation in parameter									
parameter	-5%	-4%	-3%	-2%	-1%	+1%	+2%	+3%	+4%	+5%
cond	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
eta	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
finlet	7.07	-5.68	-4.28	-2.86	-1.44	1.45	2.91	4.39	5.88	7.38
gama	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
qliht	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
qrad	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
qradv	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
qvapl	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
rodqq	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
sdent	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
sent	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
sphts	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
tmasg	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
tmasl	0.001	-0.001	-0.001	0.0	0.0	0.001	0.002	0.003	0.003	0.004
tmasv	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
tmome	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
tmoml	0.003	0.003	0.003	0.002	0.001	-0.001	-0.002	-0.003	-0.004	-0.006
tmomv	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
tnrgl	0.001	0.001	0.001	0.001	0.0	0.0	0.0	0.0	0.0	0.0
tnrgv	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
xk	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
xkes	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
xkge	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
xkl	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
xkle	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
xkvls	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
xkwew	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
xkwlw	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
xkwvw	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0

relationship between each of them and **TotalPressure**. Parameters **k_tmasl**, **k_tmoml**, and **k_tnrgl** induce small changes in the response, and the remainder of the parameters have zero effect. These sensitivity results make physical sense based on assessment of the test problem. The total pressure drop should depend strongly on the loss coefficient (**k_cd**) and wall friction (**k_xkw1x**) in the dominant flow direction, with a minor dependency on the turbulent mixing between channels (**k_tmasl**, **k_tmoml**, and **k_tnrgl**).

Listing 7.1: Dakota input file for centered parameter study, with $\pm 5\%$ variation in each of 33 parameters.

```

1 strategy
   single_method
3   tabular_graphics_data
   tabular_graphics_file 'dakota_centered_33.dat'
5
method
7   # 11 total evaluations over range [0.95,1.05]
   centered_parameter_study
9   step_vector 0.01 0.01 0.01 0.01 0.01 0.01 0.01 0.01 0.01 0.01 0.01 0.01 0.01 0.01
       0.01 0.01 0.01 0.01 0.01 0.01 0.01 0.01 0.01 0.01 0.01 0.01 0.01 0.01
       0.01 0.01 0.01 0.01 0.01
   steps_per_variable 5
11
variables
13   continuous_design      33
   initial_point          1.0 1.0 1.0 1.0 1.0 1.0 1.0 1.0 1.0 1.0 1.0 1.0 1.0 1.0
       1.0 1.0 1.0 1.0 1.0 1.0 1.0 1.0 1.0 1.0 1.0 1.0 1.0 1.0 1.0 1.0 1.0
15   descriptors            'k_eta' 'k_gama' 'k_sent' 'k_sdent' 'k_tmasv' 'k_tmasl' '
       k_tmasg' 'k_tmomv' 'k_tmome' 'k_tmoml' 'k_xk' 'k_xkes' 'k_xkge' 'k_xkl' '
       k_xkle' 'k_xkvls' 'k_xkwvw' 'k_xkwlw' 'k_xkwew' 'k_qvapl' 'k_tnrgv' 'k_tnrgl'
       'k_rodqq' 'k_qradd' 'k_qradv' 'k_qliht' 'k_sphts' 'k_cond' 'k_xkwvx' '
       k_xkwlx' 'k_cd' 'k_cdfb' 'k_wkr'
17 interface
   fork
19   asynchronous
       evaluation_concurrency = 60
21   analysis_driver = 'dakota-vera-analysis'
   # extract TotalPressure metric (length 1)
23   analysis_components = 'TotPress'
   parameters_file = 'params.in'
25 #   aprepro
       results_file = 'results.out'
27   failure_capture recover NaN
   work_directory
29   directory_tag
   named 'workdir'
31   file_save directory_save

33 responses
   num_response_functions = 1
35   descriptors
       'TotalPressure'
37   no_gradients
   no_hessians

```

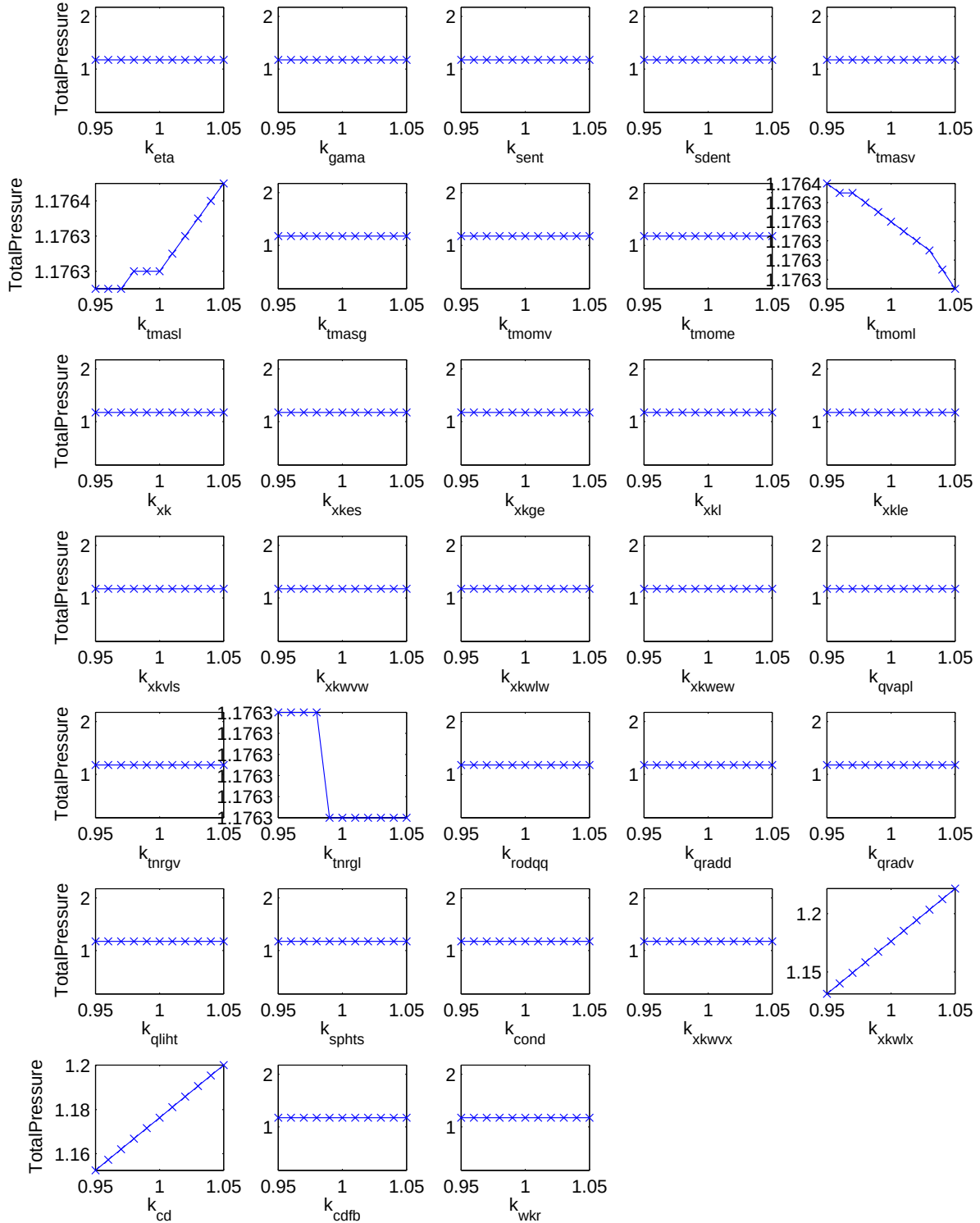


Figure 7.1: Sensitivities of TotalPressure to each of 33 variables varied over $\pm 5\%$ in a centered parameter study. Note that most parameters had identically zero variation and **k_tmasl**, **k_tmoml**, and **k_tnrgl** have only slight variation.

7.2.2 Latin hypercube sampling studies

Latin hypercube sampling for sensitivity analysis is described in Section 3.2.3. Since all the COBRA-TF parameters affect model form in the solution, they are taken to have uniform distributions on the interval $[-10\%, 10\%]$. Upon conducting initial LHS studies, the simulation failed to converge when run with deviation of `k_tmas1` $> 5\%$. Two possibilities for this failure are likely and would require additional investigation to fully diagnose. One possibility is that perturbing this parameter along with others pertinent to overall mass conservation introduces an inherent inconsistency precluding steady mass conservation. Another possibility involves the temporal stability of running COBRA-TF to a steady state. Increasing `k_tmas1` too much may violate the stability requirements of the COBRA-TF time integration method such that a steady-state is not reached. We did not pursue this further but instead decreased the range of perturbation of this parameter such that converged steady-state solutions were obtained in a stable manner. Therefore a uniform distribution on $[-5\%, 5\%]$ was used for this parameter.

Dakota Input: To change the Dakota input from a centered parameter study to a LHS study, the method and variables specifications change. The method block now prescribes a Latin hypercube sampling study. The number of samples is specified to be $N = 10 \times (M =) 33$ parameters, or $N = 330$:

```
method
  sampling
    sample_type lhs
    samples = 330
    seed = 52983
```

The variables section changes to use uncertain variables with a uniform distribution on $[0.9, 1.1]$, except for `k_tmas1` $\in [0.95, 1.05]$:

```
uniform_uncertain = 33
upper_bounds      1.1 1.1 1.1 1.1 1.1 1.05 1.1 1.1 1.1 1.1 1.1 1.1 1.1
                  1.1 1.1 1.1 1.1 1.1 1.1 1.1 1.1 1.1 1.1 1.1 1.1
                  1.1 1.1 1.1 1.1 1.1 1.1 1.1
lower_bounds      0.9 0.9 0.9 0.9 0.9 0.95 0.9 0.9 0.9 0.9 0.9 0.9 0.9
                  0.9 0.9 0.9 0.9 0.9 0.9 0.9 0.9 0.9 0.9 0.9 0.9
                  0.9 0.9 0.9 0.9 0.9 0.9 0.9
```

The full Dakota input file `dakota_lhs_33.in` is shown in Listing 7.2.

Results and Discussion: Relevant outputs generated by Dakota include the correlation coefficients in the screen output `dakota_lhs_33.out`, and tabulated data in `dakota_lhs_33.dat`. The portion of the Dakota console output with the partial correlation coefficients is shown in Figure 7.2. The output is easier to comprehend when plotted with Matlab, as shown in the bar graph in Figure 7.3. The partial correlation coefficients near 1.0 indicate that there is a strong linear correlation between `k_cd` and `TotalPressure` and between `k_xkw1x` and `TotalPressure`, consistent with physical intuition that total pressure drop should depend linearly on both the axial grid spacer loss coefficient and the axial wall friction coefficient. No other parameters are strongly significant by this measure (greater than 0.5), although several are greater than 0.1 and could be considered for inclusion in follow-on analyses based on assessment of their interaction with other parameters or nonlinear behavior.

Listing 7.2: Dakota input file for Latin hypercube sampling-based sensitivity analysis study with 330 samples and uniform input distributions.

```

strategy
2   single_method
   tabular_graphics_data
4   tabular_graphics_file 'dakota_lhs_33.dat'

6 method
   sampling
8   sample_type lhs
   samples = 330
10  seed = 52983

12 variables
   uniform_uncertain = 33
14  upper_bounds      1.1 1.1 1.1 1.1 1.1 1.05 1.1 1.1 1.1 1.1 1.1 1.1 1.1 1.1 1.1
      1.1 1.1 1.1 1.1 1.1 1.1 1.1 1.1 1.1 1.1 1.1 1.1 1.1 1.1 1.1 1.1 1.1 1.1
   lower_bounds      0.9 0.9 0.9 0.9 0.9 0.95 0.9 0.9 0.9 0.9 0.9 0.9 0.9 0.9 0.9
      0.9 0.9 0.9 0.9 0.9 0.9 0.9 0.9 0.9 0.9 0.9 0.9 0.9 0.9 0.9 0.9 0.9 0.9
16  descriptors      'k_eta' 'k_gama' 'k_sent' 'k_sdent' 'k_tmasv' 'k_tmasl' '
      k_tmasg' 'k_tmomv' 'k_tmome' 'k_tmoml' 'k_xk' 'k_xkes' 'k_xkge' 'k_xkl' '
      k_xkle' 'k_xkvls' 'k_xkwvw' 'k_xkwlw' 'k_xkwew' 'k_qvapl' 'k_tnrgv' 'k_tnrgl'
      'k_rodqq' 'k_qradd' 'k_qradv' 'k_qliht' 'k_sphts' 'k_cond' 'k_xkwvx' '
      k_xkwlx' 'k_cd' 'k_cdfb' 'k_wkr'

18 interface
   fork
20   asynchronous
      evaluation_concurrency = 60
22   analysis_driver = 'dakota-vera-analysis'
      # extract TotalPressure metric (length 1)
24   analysis_components = 'TotPress'
      parameters_file = 'params.in'
26 #   aprepro
      results_file = 'results.out'
28   failure_capture recover NaN
      work_directory
30   directory_tag
      named 'workdir'
32   file_save directory_save

34 responses
   num_response_functions = 1
36   descriptors
      'TotalPressure'
38   no_gradients
   no_hessians

```

Partial Correlation Matrix between input and output:

```
      TotalPressure
k_eta  7.48866e-02
k_gama -3.07977e-02
k_sent -3.36640e-02
k_sdent -6.83465e-02
k_tmasv -3.48366e-02
k_tmasl  1.11817e-01
k_tmasg -1.88872e-01
k_tmomv -1.15407e-01
k_tmome  1.68634e-02
k_tmoml  1.90197e-02
k_xk  8.00607e-02
k_xkes -4.52322e-02
k_xkge -6.88362e-02
k_xkl  3.59859e-02
k_xkle -3.31872e-02
k_xkvls  1.08367e-01
k_xkwvw -9.66830e-02
k_xkwlw  1.43162e-01
k_xkwew -7.26037e-03
k_qvap1 -8.55824e-02
k_tnrgv -2.87188e-02
k_tnrgl -5.43819e-03
k_rodqq  1.81288e-02
k_qradd -2.23918e-02
k_qradv -1.10458e-02
k_qliht -8.77747e-03
k_sphts -4.76993e-02
k_cond -3.93625e-02
k_xkwvx  3.31933e-02
k_xkwlx  9.99585e-01
k_cd  9.98520e-01
k_cdfb -1.86595e-02
k_wkr  2.35471e-02
```

Figure 7.2: Dakota console output showing partial correlations for the COBRA-TF simulated Progression Problem 6.

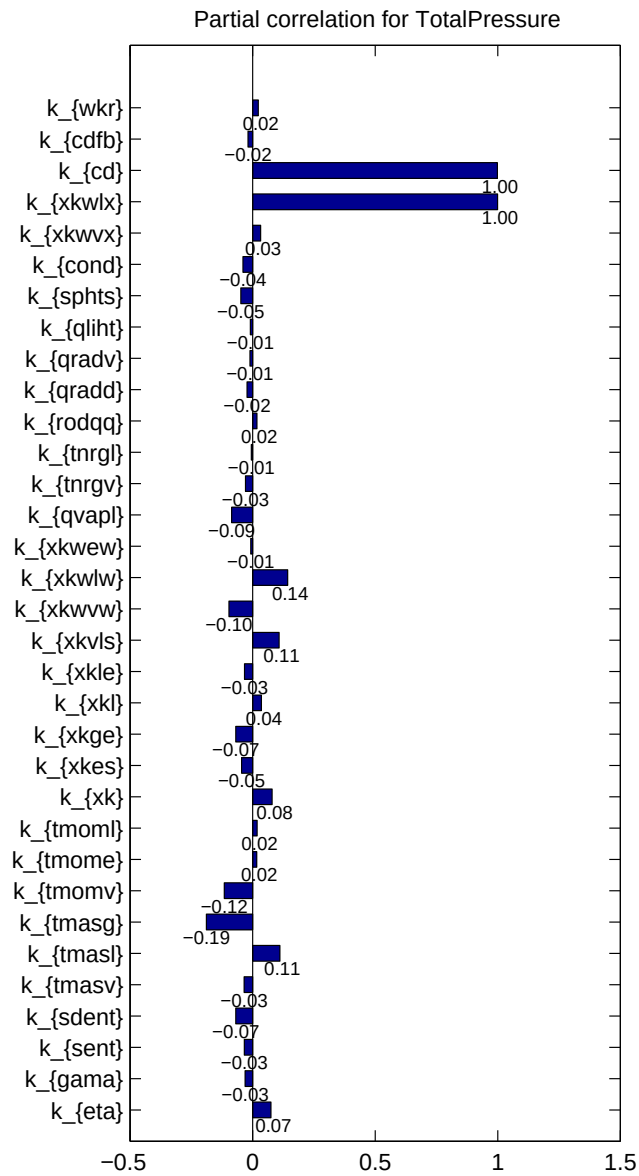


Figure 7.3: Bar graph showing partial correlation for each of 33 variables with **TotalPressure** for the COBRA-TF simulated Progression Problem 6.

Figure 7.4 displays scatter plots generated with Matlab from the Dakota-generated tabular data file. Each plot shows the overall relationship between each parameter and `TotalPressure`, with the additional vertical variation being due to the other parameters not plotted. The green lines are a linear regression on the displayed data, indicating the strength of the linear parameter response relationship, again strongest for `k_cd` and `k_xkw1x`. There is no distinguishable input/output trend for the other variables, and no patterns in the scatter cloud to suggest concern about strong nonlinear or interaction effects.

7.2.3 Morris Screening

Dakota Input: To change the study from LHS to the Morris screening method described in Section 3.2.4, one need only change the Dakota method specification to the following:

```
method
  psuade_moat
  partitions = 9    #to generate 10 levels
  samples = 340    # must be integer multiple of (num_vars + 1)
  seed = 20
```

The full Dakota input file is available in `examples/CobraTF/MorrisStudies/dakota_morris_33.in`.

Results and Discussion: The results of the Morris study are the modified means and standard deviations of the elementary effects in the Dakota console output, excerpted from `dakota_morris_33.out` into Figure 7.5. These indicate that inputs 30 and 31 (`k_cd` and `k_xkw1x`) have a strong main effect and some interaction effect. They also indicate that no other parameters have any effect, save 6, 10, and 22, corresponding to `k_tmas1`, `k_tmom1`, and `k_tnrg1`, which each have very small effects. This is physically consistent. The main forces impacting pressure drop are wall friction and loss coefficients with minor impacts due to turbulent mixing.

7.2.4 Screening to Reduce Parameters

Figure 7.6 summarizes the results of the sensitivity studies conducted, showing correlations from the LHS study, effects from the Morris study, and variation seen in the centered parameter study. The various sensitivity methods are consistent with each other. Based on these results, the following studies will use only parameters inducing variation in `TotalPressure`: `k_tmas1`, `k_tmom1`, `k_tnrg1`, `k_xkw1x`, and `k_cd`. Of these, only the last two have strong effects. Because the axial grid spacer loss coefficient and the axial wall friction coefficient appear in the axial momentum equations as multipliers of the axial velocity squared, the total pressure drop can only be related to the sum of these two coefficients. The impact of this sort of parameter nonidentifiability on model calibration will be illustrated in the analyses of Section 7.3.3.

LHS studies with 50 and 200 samples were then conducted over the five most sensitive parameters to (1) alleviate any potential confounding from the other parameters in the sensitivity metrics and (2) generate simulation data to use for calibration studies in the next section. The input and output files for these studies are omitted from the text (as the sensitivity results are similar), but are available in `examples/CobraTF/LHSStudies/`:

```
dakota_lhs_5.200.dat  dakota_lhs_5.200.out  dakota_lhs_5.50.in
dakota_lhs_5.200.in   dakota_lhs_5.50.dat  dakota_lhs_5.50.out
```

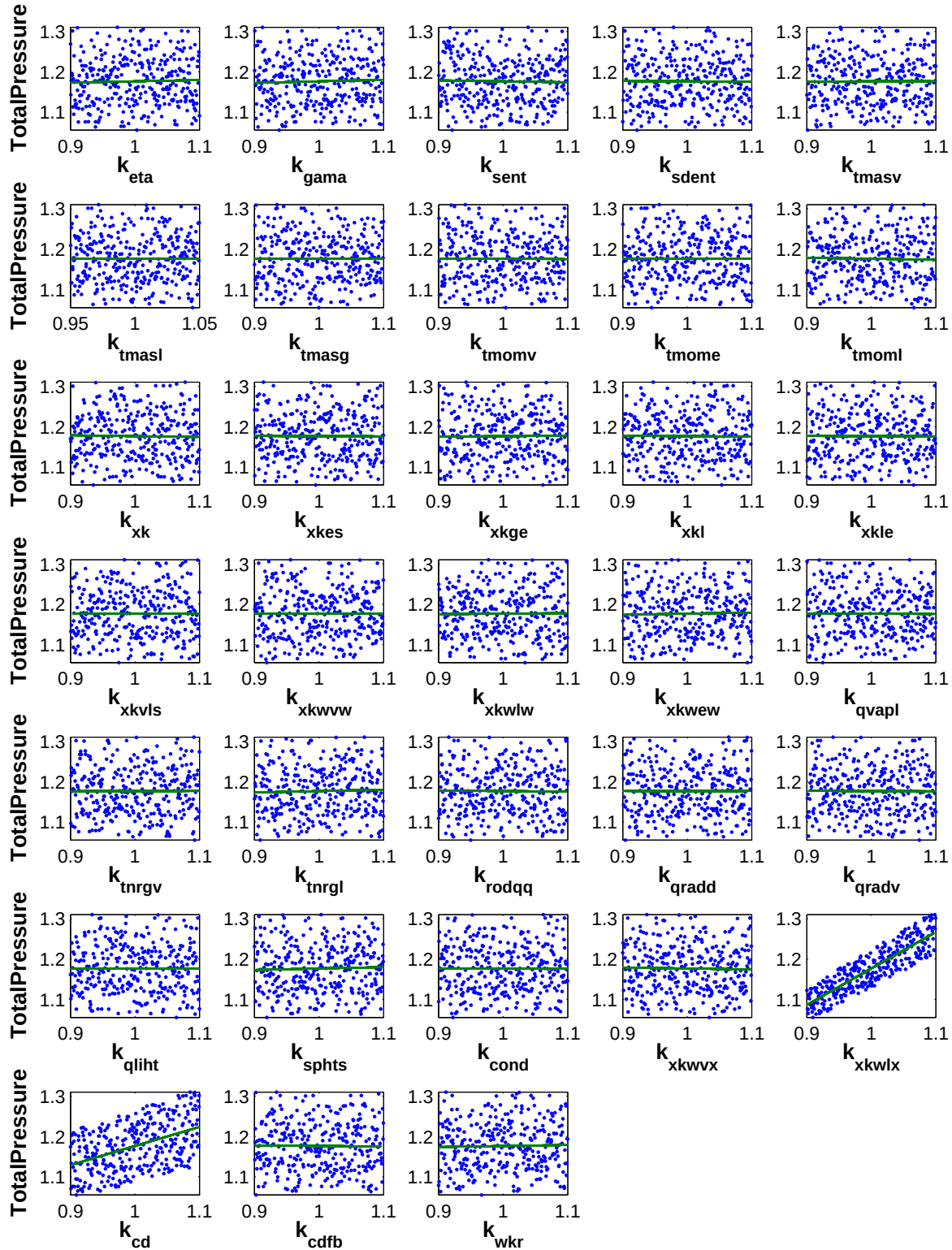


Figure 7.4: Scatter plots with regression lines for each of 33 variables with **TotalPressure** for the COBRA-TF simulated Progression Problem 6.

```

>>>>> PSUADE MOAT output for function 0:

*****
***** MOAT Analysis *****
-----
Input   1 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input   2 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input   3 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input   4 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input   5 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input   6 (mod. mean & std) =  6.4800e-05  2.2768e-05
Input   7 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input   8 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input   9 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  10 (mod. mean & std) =  2.2320e-04  1.3041e-04
Input  11 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  12 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  13 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  14 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  15 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  16 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  17 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  18 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  19 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  20 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  21 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  22 (mod. mean & std) =  9.0000e-06  9.4868e-06
Input  23 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  24 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  25 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  26 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  27 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  28 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  29 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  30 (mod. mean & std) =  1.8036e-01  7.0671e-03
Input  31 (mod. mean & std) =  9.5857e-02  7.8750e-03
Input  32 (mod. mean & std) =  0.0000e+00  0.0000e+00
Input  33 (mod. mean & std) =  0.0000e+00  0.0000e+00
<<<<< Function evaluation summary: 340 total (340 new, 0 duplicate)

<<<<< Iterator psuade_moat completed.
<<<<< Single Method Strategy completed.
DAKOTA execution time in seconds:
  Total CPU      =      1.28

```

Figure 7.5: Dakota output showing modified means and standard deviations of elementary effects for the COBRA-TF simulated Progression Problem 6.

parameter	partial correlation	simple correlation	morris main	morris interaction	CPS variation
k_eta	0.07	0.03			
k_gama	-0.03	0.04			
k_sent	-0.03	-0.02			
k_sdent	-0.07	-0.01			
k_tmasv	-0.03	0.00			
k_tmasl	0.11	0.00	6.48E-05	2.28E-05	medium
k_tmasg	-0.19	-0.01			
k_tmomv	-0.12	-0.01			
k_tmome	0.02	0.00			
k_tmoml	0.02	-0.02	2.23E-04	1.30E-04	medium
k_xk	0.08	-0.02			
k_xkes	-0.05	0.00			
k_xkge	-0.07	0.01			
k_xkl	0.04	-0.01			
k_xkle	-0.03	0.00			
k_xkvl	0.11	-0.01			
k_xkwvw	-0.10	0.01			
k_xkwlw	0.14	0.01			
k_xkwew	-0.01	0.03			
k_qvap	-0.09	-0.01			
k_tnrgv	-0.03	0.00			
k_tnrgl	-0.01	0.03	9.00E-06	9.49E-06	low
k_rodqq	0.02	-0.01			
k_qradd	-0.02	0.00			
k_qradv	-0.01	0.00			
k_qliht	-0.01	0.00			
k_sphts	-0.05	0.03			
k_cond	-0.04	0.00			
k_xkwvx	0.03	-0.02			
k_xkwlx	1.00	0.88	1.80E-01	7.07E-03	high
k_cd	1.00	0.46	9.59E-02	7.88E-03	high
k_cdfb	-0.02	-0.01			
k_wkr	0.02	0.02			

Figure 7.6: Summary of sensitivity analysis results for 33 COBRA-TF parameters. Missing values are identically zero. Highlighted rows will be used in subsequent studies.

In particular the tabular data files are used in the follow-on Bayesian calibration and surrogate generation examples.

7.3 Calibration Studies

In this section we emulate the model calibration processes described in Sections 5.1.1 and 6.2 by generating synthetic data and then applying Dakota algorithms to determine the parameter values yielding the best match between model and data. The 10 synthetic data points, generated by adding independent and identically distributed Gaussian noise $\varepsilon \sim \mathcal{N}(0, \sigma^2)$ for $\sigma = 0.025$ to the

nominal TotalPressure = 1.17304 bar, are placed in `ctf_dat.txt` for consumption by Dakota. They are:

```
1.21626870795683
1.16149729403381
1.1866271753701
1.17683037400566
1.17908641397936
1.12766758338348
1.17093102211906
1.2072835297607
1.18316329384675
1.19396014684855
```

7.3.1 Deterministic Calibration

Dakota Input: A local gradient-based algorithm should perform well for this model calibration problem as the global sensitivity analysis above revealed smooth and linear trends. The Dakota input file for this problem is shown in Listing 7.3. Highlights include: use of the NL2SOL method (line 7) for local calibration; use of design variables (line 11) instead of uncertain variables, since we are calibrating; use of calibration terms in the responses section (line 35); and input of the external data file into Dakota (line 38) for computing the least-squares residuals.

We also repeat the same study, calibrating only the two most influential variables (holding the other three fixed at nominal) in hopes of improving the solution by eliminating unidentifiable (insensitive) parameters. This requires changing the variables section to:

```
continuous_design = 2
initial_point      1.0103  0.9799
upper_bounds       1.05   1.1
lower_bounds       0.95   0.9
descriptors        'k_tmasl' 'k_cd'
```

The complete input file is available in `examples/CobraTF/Calibration/dakota_calibration.2.in`.

Results and Discussion: For each study, the end of the Dakota output (in files `examples/CobraTF/Calibration/dakota_calibration_[25].out`) indicates the best solution found, as shown in Figure 7.7. The best values of the parameters are shown, followed by the residuals between the model calculations and data. A summary of the study initial and final values/residual norms appears in Table 7.3. In both studies, the optimization solver made progress toward recovering the correct nominal values of `k_xkw1x` and `k_cd`, but no appreciable progress on the other parameters. This is expected as the TotalPressure response is not sensitive to these input parameters, rendering them essentially unidentifiable from the provided data.

7.3.2 Surrogate Construction

The analysis of Section 7.2 resulted in a reduction of the initial 33 COBRA-TF parameters to a final set of five sensitive parameters with respect to induced variation in total pressure drop: `k_tmom1`, `k_tnrgr1`, `k_xkw1x`, and `k_cd`.

Listing 7.3: Dakota input file for deterministic local gradient-based calibration of five key parameters in the COBRA-TF simulated Progression Problem 6.

```

1 strategy
  single_method
3  tabular_graphics_data
  tabular_graphics_file 'dakota_calibration_5.dat'
5
method
7  nl2sol
  convergence_tolerance 1.0e-6
9
variables
11 continuous_design = 5
  initial_point      1.0103      0.9766      0.9888      0.9773      0.9799
13 upper_bounds      1.05 1.1 1.1 1.1 1.1
  lower_bounds      0.95 0.9 0.9 0.9 0.9
15 descriptors      'k_tmasl' 'k_tmoml' 'k_tnrgl' 'k_xkwlx' 'k_cd'
17 interface
  fork
19    asynchronous
    # use max concurrency = 11
21    # evaluation_concurrency = 60
    analysis_driver = 'dakota-vera-analysis'
23    # extract TotalPressure metric (length 1)
    analysis_components = 'TotPress'
25    parameters_file = 'params.in'
  #
27    aprepro
    results_file      = 'results.out'
    failure_capture recover NaN
29    work_directory
    directory_tag
31    named 'workdir'
    file_save directory_save
33
responses
35  calibration_terms = 1
    descriptors
37    'TotalPressure'
    calibration_data_file = 'ctf_dat.txt'
39    freeform
    num_experiments = 1
41    num_replicates = 10
  numerical_gradients
43    central
    # coarse FD step as cobra might not be sensitive enough
45    fd_step_size = 1.0e-2
  no_hessians

```

```

<<<<< Function evaluation summary: 48 total (38 new, 10 duplicate)
<<<<< Best parameters      =
      1.0500000000e+00 k_tmasl
      9.0000000000e-01 k_tmoml
      9.8880000000e-01 k_tnrgl
      9.9500856366e-01 k_xkwlx
      1.0177721341e+00 k_cd
<<<<< Best residual norm =  7.4038381819e-02; 0.5 * norm^2 =  2.7408409912e-03
<<<<< Best residual terms  =
      -3.5938707957e-02
       1.8832705966e-02
      -6.2971753701e-03
       3.4996259943e-03
       1.2435860206e-03
       5.2662416617e-02
       9.3989778809e-03
      -2.6953529761e-02
      -2.8332938467e-03
      -1.3630146849e-02
<<<<< Best constraint values =
<<<<< Best data not found in evaluation cache
Confidence Interval for k_tmasl is [ -6.8687978090e+01,  7.0787978090e+01 ]
Confidence Interval for k_tmoml is [  9.0000000000e-01,  9.0000000000e-01 ]
Confidence Interval for k_tnrgl is [  9.8880000000e-01,  9.8880000000e-01 ]
Confidence Interval for k_xkwlx is [  9.9500856366e-01,  9.9500856366e-01 ]
Confidence Interval for k_cd is [  1.0177721341e+00,  1.0177721341e+00 ]

```

Figure 7.7: Dakota console output showing final results for calibration with five parameters.

Table 7.3: Summary of five- and two-parameter calibration for the COBRA-TF simulated Progression Problem 6, initial and final parameter values and residual norm.

	5 parameter case		2 parameter case	
	initial	final	initial	final
k_tmasl	1.0103	1.0500	1.0103	1.0414
k_tmoml	0.9766	0.9000	1.0	1.0
k_tnrgl	0.9888	0.9888	1.0	1.0
k_xkwlx	0.9773	0.9950	1.0	1.0
k_cd	0.9799	1.0178	0.9799	1.0083
$\frac{1}{2} r ^2$	0.0085	0.0027	0.0085	0.0027

Dakota Input: Listing 7.4 shows the Dakota input file `ctf_gp_eval.in` for fitting a kriging model with constant trend to total pressure drop as a function of these five parameters, and using this fit to predict total pressure drop on a set of 20 validation runs. The kriging fit itself was based

on the results of using Dakota to run COBRA-TF on a Latin hypercube sample of size $N = 50$. The output of these COBRA-TF runs was written to the file `dakota.pstudy.dat`, which is subsequently read into this job via the `import_points` option (line 35).

Results and Discussion: Figure 7.8 compares the emulator predictions with the COBRA-TF calculations of total pressure drop for the 20 validation runs. The left panel plots predicted value against calculated value, with the resulting points falling very close to the desired 45° line. The right panel plots the standardized residuals (calculated minus predicted total pressure drop divided by standard error of predicted total pressure drop) against calculated value, indicating a fairly constant scatter in the standardized residuals around the zero line across the spectrum of calculated values. This result is desirable in that it both indicates unbiased prediction and supports the assumption of a homogeneous (constant) process variance σ^2 . Furthermore, we can apply a simple screening procedure for outliers to this standardized residual plot. First compute the interquartile range (IQR) of the standardized residuals, defined as the 75-th percentile minus the 25-th percentile of the standardized residuals. Second compute lower and upper bounds as the 25-th percentile minus $1.5 \times \text{IQR}$ and the 75-th percentile plus $1.5 \times \text{IQR}$, respectively. Any standardized residuals falling outside these bounds are flagged for further investigation as potential outliers. For this kriging fit, the lower and upper bounds calculated in this way were $(-2.40, 3.38)$, while the range of the standardized residuals themselves was well within these bounds at $(-1.33, 2.38)$.

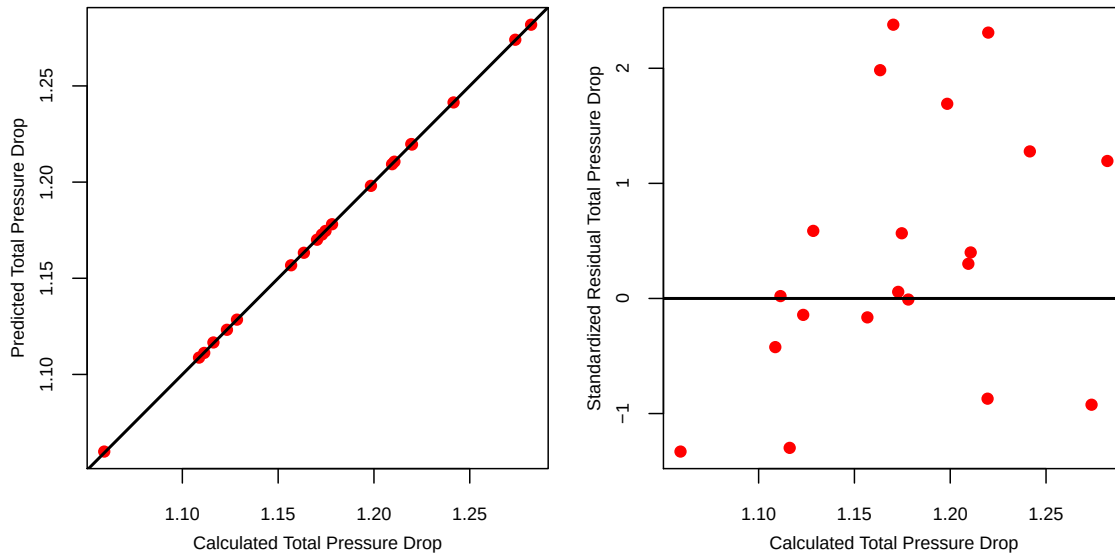


Figure 7.8: Predicted vs. calculated total pressure drop (left panel) and standardized residual vs. calculated total pressure drop (right panel) for 20 validation runs (red circles).

The leave-one-out cross-validation RMSPE of this kriging emulator, requested with the option `press` (line 34), evaluated to $3.2415584202\text{e-}04$. When validation samples are available, as in this example, they can be used to compute a validation RMSPE. For N_v validation samples, calculate

$$RMSPE = \sqrt{\frac{1}{N_v} \sum_{i=1}^{N_v} (C_i - P_i)^2},$$

Listing 7.4: Dakota input file producing predictions for 20 validation runs from a GP emulator with estimated constant trend for the COBRA-TF simulated Progression Problem 6.

```

# Build and evaluate a Gaussian process emulator of COBRA-TF output
2 # at a user specified set of points

4 strategy
    single_method
6     method_pointer = 'EvalSurrogate'
    tabular_graphics_data
8     tabular_graphics_file = 'ctf_gp_evals.dat'

10 # Method to perform evaluations of the emulator

12 method
    id_method = 'EvalSurrogate'
14    model_pointer = 'SurrogateModel'

16    # Verbose will show the type form of the surrogate model
    output verbose

18
    sampling
20    sample_type lhs
    samples = 20
22    seed = 150

24 # Surrogate model specification
model
26    id_model = 'SurrogateModel'
    surrogate global
28    # GP model
    gaussian_process surfpack
30    trend
        constant
32    # compute and print diagnostics after build
    metrics 'rsquared' 'root_mean_squared'
34    press
    import_points = 'dakota_pstudy.dat' annotated
36

variables,
38    uniform_uncertain =    5
    upper_bounds          1.05 1.1 1.1 1.1 1.1
40    lower_bounds          0.95 0.9 0.9 0.9 0.9
    descriptors            'k_tmasl' 'k_tmoml' 'k_tnrgl' 'k_xkwlx' 'k_cd'
42

responses
44    response_functions = 1
    descriptors = 'TotalPressure'
46    no_gradients
    no_hessians

```

where P_i and C_i denote the i -th predicted and calculated values, respectively. For our validation sample of $N_v = 20$, the RMSPE evaluates to $2.508674\text{e-}04$. These RMSPE values are 0.15% and 0.11% of the observed range in the 20 calculated total pressure drops, indicating the kriging emulator possesses high accuracy for this application.

The assumption of Gaussian errors can be checked by examining normal probability plots of the standardized residuals from both leave-one-out cross validation and out-of-sample validation. Figure 7.9 shows normal probability plots for both of these cases. Theoretical quantiles from the standard normal distribution are plotted on the x -axis, while the corresponding sample quantiles of the standardized residuals are plotted on the y -axis. Consistency with the Gaussian error assumption is indicated by the plotted points exhibiting strong linear association. The simple correlation coefficients for the two cases are 0.992 and 0.965, respectively, suggesting that the Gaussian error assumption is reasonable. The standardized residuals resulting from leave-one-out cross validation are correlated, making the normal probability plot somewhat harder to interpret as the theoretical quantiles are calculated assuming independence. On the other hand, for the validation runs, the residuals were transformed to obtain uncorrelated standardized residuals for use in the normal probability plot [6].

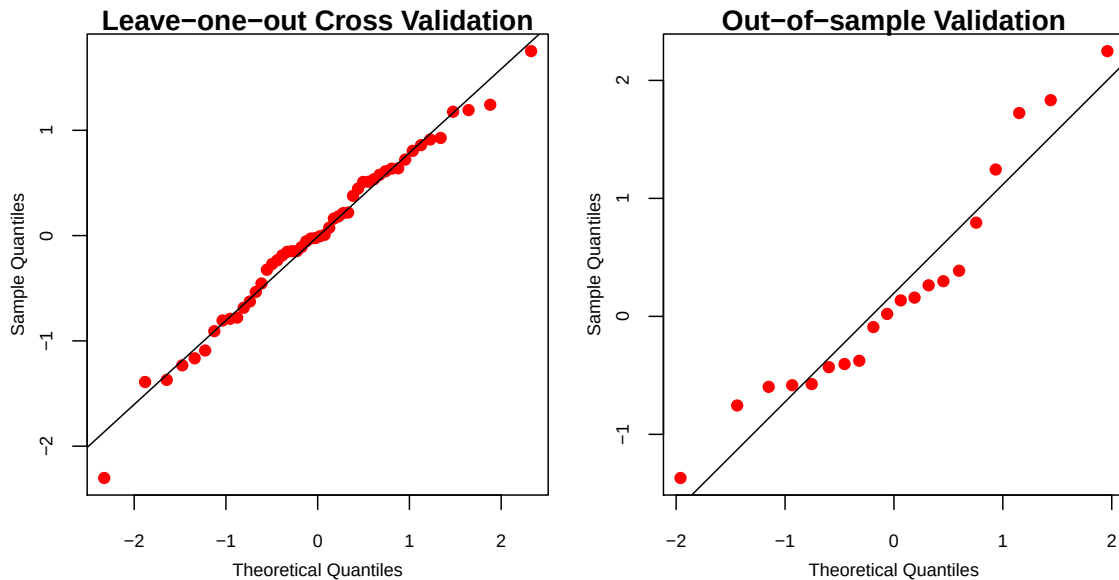


Figure 7.9: Normal probability plots based on standardized residuals from leave-one-out cross validation (left panel) and out-of-sample validation (right panel).

7.3.3 Bayesian Calibration

We illustrate here the use and verification of the Bayesian model calibration techniques, discussed in Section 6.2, for COBRA-TF. We focus on the parameters `k_tmas1`, `k_tmom1`, `k_tnrg1`, `k_xkw1x`, and `k_cd` which, as detailed in the sensitivity analysis of Section 7.2, had non-negligible values for the Morris elementary effect statistics. The synthetic data of Section 7.3, generated as detailed in Section 6.2.5, is consistent with the likelihoods employed in DRAM, DREAM, and GPMSA.

Analysis reveals that the total pressure drop can be accurately quantified by a linear relation

between `k_xkwlx` and `k_cd`. This is important for three reasons.

- This causes `k_xkwlx` and `k_cd` to be nonidentifiable in the sense that the input-to-output relation is not unique – see Chapter 6 of [34] for details regarding the ramifications of non-identifiable parameter sets. However, by specifying prior bounds on parameters, we can ensure that we obtain a proper posterior density.
- We can show that the variable

$$\eta = 0.906 \text{ k_xkwlx} + 0.478 \text{ k_cd} \quad (7.1)$$

is approximately normally distributed with mean 1.39 and standard deviation 7.91×10^{-3} . This provides an approximate analytic solution that we can employ for verification.

- This analysis predicts that `k_xkwlx` and `k_cd` should be negatively, and approximately linearly, correlated. This information can also be used to verify sampling results.

For DRAM and DREAM implementation, we employed the surrogate, constructed using the five parameters, as detailed in Section 7.3.2. GPMSA employs a Gaussian process emulator constructed from simulations run with pre-specified parameter inputs. We note that emulator uncertainty is accounted for in GPMSA analysis whereas it is not in the Dakota-QUESO DRAM and Dakota DREAM analyses. In all three packages, we enforced the parameter bounds summarized in Table 7.4. We employed uniform densities over these ranges as prior densities for each random variable.

Table 7.4: Parameter bounds used to construct prior densities.

Descriptors	<code>k_tmasl</code>	<code>k_tmoml</code>	<code>k_tnrgl</code>	<code>k_xkwlx</code>	<code>k_cd</code>
Lower Bounds	0.95	0.9	0.9	0.9	0.9
Upper Bounds	1.05	1.1	1.1	1.1	1.1

We ran the DRAM chain for 10,000 iterations and constructed 5 DREAM chains each of length 2000. As illustrated by the representative DRAM chain for `k_tmasl`, which is plotted in Figure 7.10(a), the chain has burned in by 2000 so we employed the last 8000 elements for kernel density estimation (kde). The DRAM and DREAM densities for η , given by (7.1), are compared with the approximate analytic normal density $N(1.39, (7.91 \times 10^{-3})^2)$ in Figure 7.10(b). This provides a verification test in the manner detailed in Section 6.2.4. In Figure 7.10(c), we compare GPMSA results, obtained with fixed and inferred measurement variances σ , with the approximate analytic results. In Figure 7.10(d), we show that GPMSA, with inferred variances, accurately matches the DRAM results. In combination, this verifies the accuracy of the Bayesian calibration procedures and the surrogate model developed in Section 7.3.2.

We illustrate in Figure 7.11 the joint sample points for the negatively correlated random variables `k_xkwlx` and `k_cd`. These plots reflect the parameter limits specified in Table 7.4. The fact that DRAM, DREAM and GPMSA yield the same joint sample plots, which are consistent with the linear relation (7.1), further verifies the accuracy of all three methods for this implementation of COBRA-TF. The input decks for DRAM and DREAM are provided in Listings 7.5 and 7.6.

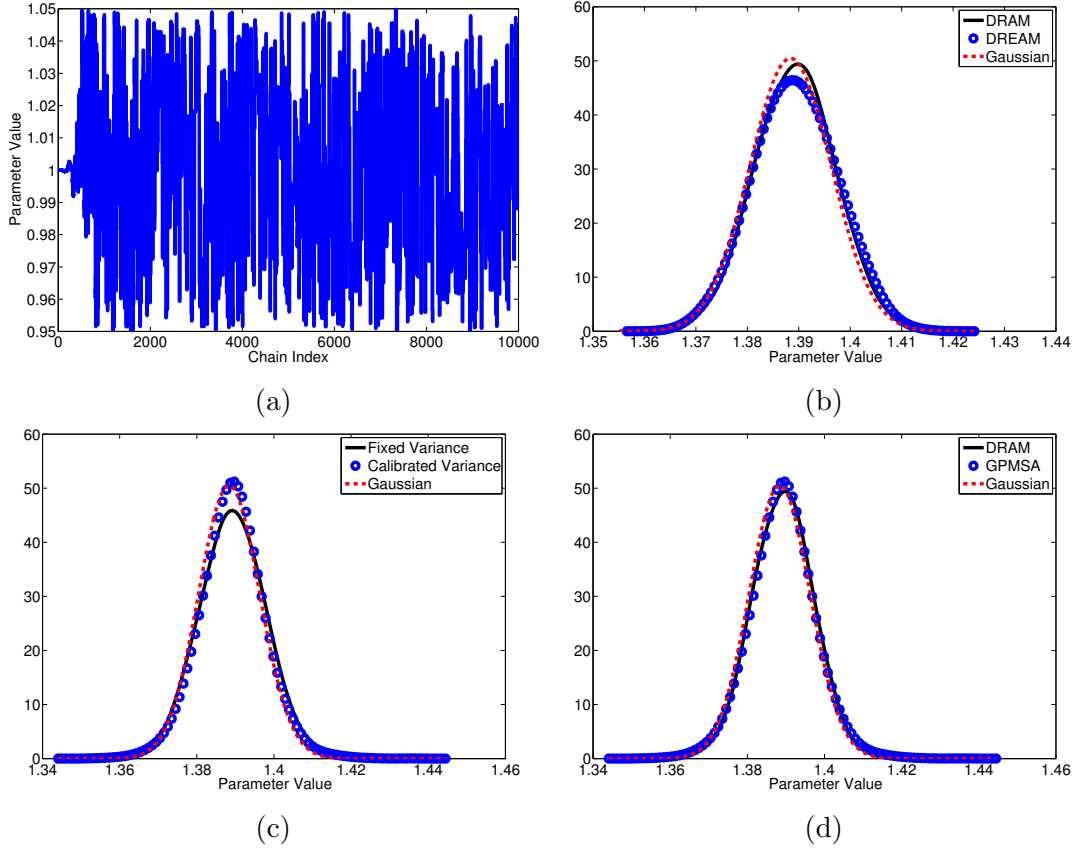


Figure 7.10: (a) Representative DRAM chain for `k_tmas1` showing burn-in by 2000. (b) Verification of the densities for η provided by DRAM and DREAM through comparison with the analytic Gaussian solution. (c) and (d) Verification of GPMSA for fixed and inferred measurement variance σ .

It was noted in Section 6.2.3 that once Bayesian chains have been constructed, values from the chain can be used as inputs to surrogate models to compute calibrated predictions for additional quantities of interest. This is illustrated in Figure 7.12, where we compare credible intervals for the total pressure drop computed using GPMSA with the surrogate model evaluated using inputs from the DRAM chain. The nominal total pressure drop value of 1.176 bar, computed using COBRA-TF with nominal parameter values, lies within one standard deviation of the mean. Furthermore, the prediction is essentially Gaussian due to the linearity of the response in this operating regime. The Dakota input deck is provided in Listing 7.7 where one can note that it inputs DRAM values from the file `ctf_gp_evals_dram.dat`.

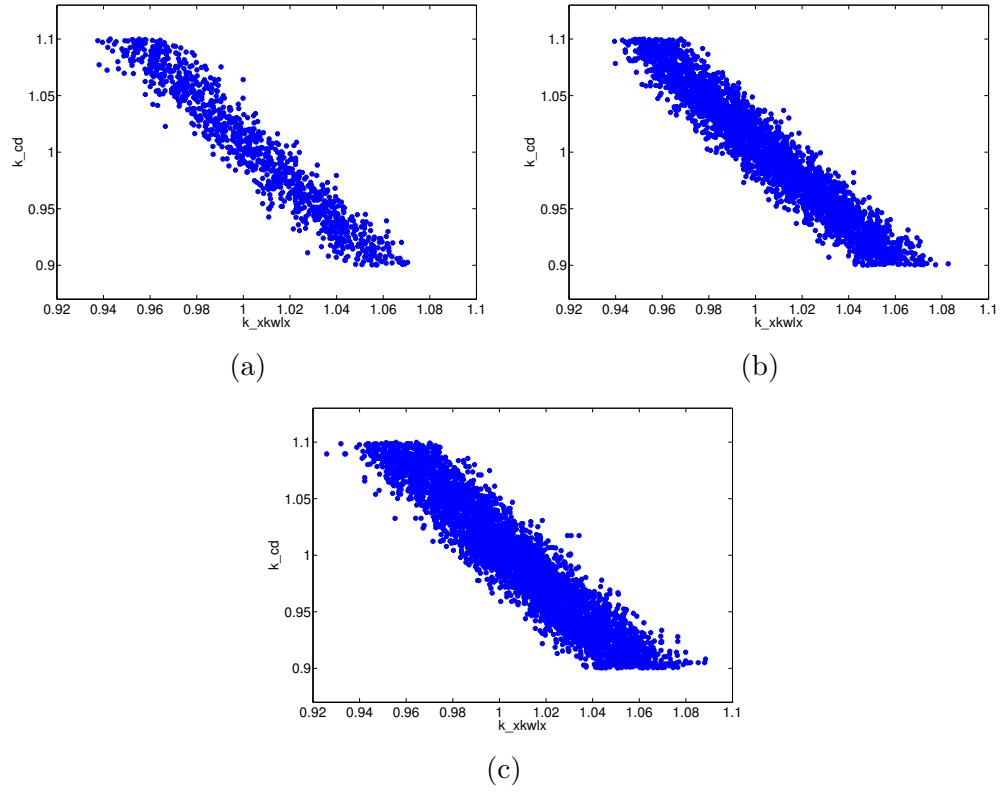


Figure 7.11: Joint sample points for k_{xkw1x} and k_{cd} provided by (a) DRAM, (b) DREAM and (c) GPMSA.

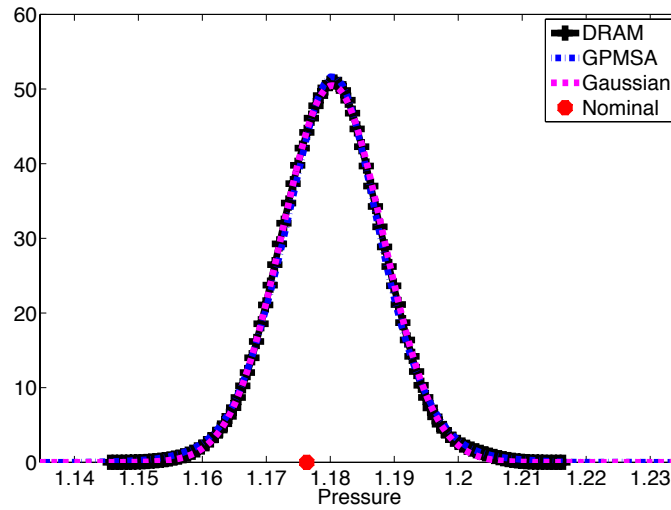


Figure 7.12: Comparison of the credible interval constructed using GPMSA and the surrogate with inputs from the DRAM chain. The nominal total pressure drop value of 1.176 bar was computed using COBRA-TF.

Listing 7.5: Input for Dakota-QUESO DRAM as applied to the COBRA-TF surrogate for Progression Problem 6.

```

1 strategy
   single_method
3
method
5   bayes_calibration queso
   emulator
7   gaussian_process surfpack
   import_points = 'dakota_5pstudy.dat' annotated
9   mcmc_type dram
   rejection delayed
11  metropolis adaptive
   samples = 10000
13  seed = 52983

15 variables
   active design
17  continuous_design = 5
   #initial_point      1.05 1.0 1.0 1.04
19  upper_bounds        1.05 1.1 1.1 1.1 1.1
   lower_bounds        0.95 0.9 0.9 0.9 0.9
21  descriptors         'k_tmasl' 'k_tmoml' 'k_tnrgl' 'k_xkwlx' 'k_cd'

23 interface
   direct
25  analysis_driver = 'text_book'

27 responses
   calibration_terms = 1
29  calibration_data_file = 'ctf_dat.txt'
   freeform
31  num_experiments = 1
   num_replicates = 10
33  num_std_deviations = 1
   descriptors
35  'TotalPressure'
   no_gradients
37  no_hessians

```

Listing 7.6: Input for Dakota DREAM as applied to the COBRA-TF surrogate for Progression Problem 6.

```

1 strategy
3   single_method

5 method
   bayes_calibration dream
7   emulator
   gaussian_process surfpack
9   import_points = 'dakota_5pstudy.dat' annotated
   #mcmc_type dram
11  #rejection delayed
   #metropolis adaptive
13  samples = 10000
   chains = 5
15  seed = 52983

17 variables
   active design
19   continuous_design = 5
   upper_bounds      1.05 1.1 1.1 1.1 1.1
21  lower_bounds      0.95 0.9 0.9 0.9 0.9
   descriptors        'k_tmasl' 'k_tmoml' 'k_tnrgl' 'k_xkwlx' 'k_cd'
23

interface
25   direct
   analysis_driver = 'text_book'
27

responses
29   calibration_terms = 1
   calibration_data_file = 'ctf_dat.txt'
31   freeform
   num_experiments = 1
33   num_replicates = 10
   num_std_deviations = 1
35   descriptors
       'TotalPressure'
37   no_gradients
   no_hessians

```

Listing 7.7: Dakota input used to construct credible interval for total pressure drop using DRAM evaluations.

```

# Build and evaluate a Gaussian process emulator of COBRA-TF output
2 # at a user specified set of points

4 strategy
  single_method
6 #environment
  method_pointer = 'EvalSurrogate'
8  tabular_graphics_data
  tabular_graphics_file = 'ctf_gp_evals_dram.dat'
10
# Method to perform evaluations of the emulator
12
method
14  id_method = 'EvalSurrogate'
  model_pointer = 'SurrogateModel'
16
  # Verbose will show the type form of the surrogate model
18  output verbose

20  list_parameter_study
  import_points = 'dram_result.txt' freeform
22
# Surrogate model specification
24 model
  id_model = 'SurrogateModel'
26  surrogate global
    # GP model
28    gaussian_process surfpack
    trend
30    constant
    # compute and print diagnostics after build
32    metrics 'rsquared' 'root_mean_squared'
    press
34    import_points = 'dakota_5pstudy.dat' annotated

36 variables,
  uniform_uncertain = 5
38  upper_bounds      1.05 1.1 1.1 1.1 1.1
  lower_bounds      0.95 0.9 0.9 0.9 0.9
40  descriptors       'k_tmasl' 'k_tmoml' 'k_tnrgl' 'k_xkwlx' 'k_cd'

42 responses
  response_functions = 1
44  descriptors = 'TotalPressure'
  no_gradients
46  no_hessians

```

Appendix A

General Linear Model Verification Test Suite

This appendix provides additional technical details for the linear model introduced in Section 2.2. Here we describe in detail a particular class of problems for which results of the QUESO DRAM or Dakota DREAM sampling algorithm can be verified against analytical solutions. Specifically, we present standard results for Bayesian analysis of the linear regression model ([11], pp. 233-265). In linear regression, a N -vector of outputs y is related linearly to functions of inputs x ,

$$y = G\beta + \epsilon,$$

where the i -th row of G contains the evaluation of these regression functions at input x_i corresponding to the i -th datum y_i , β denotes the regression coefficients, and ϵ denotes the vector of observational errors. In the following, observational errors will be assumed mean-zero Gaussian, having variance $(1/\lambda)$ and possibly correlated with parametric dependencies governed by a parameter ϕ .

Marginal posterior distributions for β , λ , and ϕ are derived analytically for three increasingly challenging verification scenarios: (i) β unknown, (λ, ϕ) fixed; (ii) (β, λ) unknown, ϕ fixed; and (iii) (β, λ, ϕ) unknown. In the third scenario, the marginal posterior distributions of β , λ , and ϕ do not belong to a standard class of probability distributions (such as Gaussian or Gamma, for example), thus requiring the use of numerical methods such as quadrature to accurately estimate their normalizing constants.

For each verification scenario, QUESO samples from the applicable marginal posterior distributions for β , λ , and ϕ will be compared with the corresponding analytical results. Convergence of posterior means and other summary statistics to corresponding parameter values assumed for data generation will also be monitored for a subset of the observational error correlation structures investigated.

Section A.1 describes how data are simulated and presents the three verification scenarios of interest with corresponding analytical results for the desired posterior distributions in each scenario. Section A.2 defines the verification tests currently implemented.

A.1 Verification Scenarios

Let ε_t denote a mean-zero Gaussian stochastic process having covariance function $c(t_1, t_2|\phi) = (1/\lambda) r(t_1, t_2|\phi)$, where $\lambda > 0$ and $r(\cdot, \cdot|\phi)$ is a correlation function for $\phi \in \Phi$. We assume the following quantities are specified:

1. Nominal parameter settings $(\beta_0, \lambda_0, \phi_0)$ with $\beta_0 \in \mathfrak{R}^{N_\beta}$, $\lambda_0 > 0$, and $\phi_0 \in \text{int}(\Phi)$
2. Indices $\{t_1, t_2, \dots, t_N\}$.

We denote the multivariate Gaussian distribution having location vector μ and covariance matrix Σ by $\mathcal{N}(\mu, \Sigma)$. For specified N , generate a N -vector of errors

$$(\varepsilon_{t_1}, \dots, \varepsilon_{t_N}) \sim \mathcal{N}(0_N, (1/\lambda_0)R(\phi_0)) ,$$

where 0_N is the N -vector of zeroes and the (i, j) element of $R(\phi_0)$ is given by $r(t_i, t_j|\phi_0)$. To complete the data generation process, we sample M -dimensional covariates $\{X_1, \dots, X_N\}$ independently from the distribution $\mathcal{N}(0_M, C)$, where C is a fixed covariance matrix. The i -th datum is calculated as $y_{t_i} = g^T(x_i)\beta_0 + \varepsilon_{t_i}$, where $g(\cdot)$ is a N_β -dimensional regression function. In vector-matrix form,

$$y = G\beta_0 + \epsilon ,$$

where

$$\begin{aligned} y &= (y_{t_1}, y_{t_2}, \dots, y_{t_N})^T , \\ G &= [g(x_1) \ g(x_2) \ \dots \ g(x_N)]^T , \text{ and} \\ \epsilon &= (\epsilon_{t_1}, \epsilon_{t_2}, \dots, \epsilon_{t_N})^T . \end{aligned}$$

To generate $\tilde{N}$ additional responses, we sample:

1. $\{X_{N+1}, \dots, X_{N+\tilde{N}}\}$ independently from the distribution $\mathcal{N}(0_M, C)$
2. $(\varepsilon_{t_{N+1}}, \dots, \varepsilon_{t_{N+\tilde{N}}})^T \sim \mathcal{N}\left(0_{\tilde{N}}, (1/\lambda_0)\left(\tilde{R}(\phi_0) - \bar{R}(\phi_0)R^{-1}(\phi_0)\bar{R}^T(\phi_0)\right)\right) ,$

where the (i, j) element of $\tilde{R}(\phi_0)$ is given by $r(t_{N+i}, t_{N+j}|\phi_0)$, and the (i, j) element of $\bar{R}(\phi_0)$ is given by $r(t_{N+i}, t_j|\phi_0)$. Then $y_{t_{N+i}} = g^T(x_{N+i})\beta_0 + \varepsilon_{t_{N+i}}$ and the matrix-vector form of the augmented data set follows. Note that this process of conditionally sampling errors preserves the correct joint distribution of the augmented data vector.

In the following calculations, we assume the regression matrix G is fixed, so this will not be explicitly denoted in the notation. Although (β, λ, ϕ) are fixed at $(\beta_0, \lambda_0, \phi_0)$ to generate data as above, our statistical analyses will assume some or all of $(\beta_0, \lambda_0, \phi_0)$ are unknown. The sampling distribution $f(Y|\beta, \lambda, \phi)$ of a random data set Y of size N is $\mathcal{N}(G\beta, (1/\lambda)R(\phi))$.

(Case 1.) The first verification scenario assumes that λ and ϕ are fixed at λ_0 and ϕ_0 , respectively. The Bayesian analysis places a prior distribution on β , which is given by the following:

1. $\pi(\beta) = \int_{\mathfrak{R}^{N_\beta}} \pi(\beta|\mu) \pi(\mu) d\mu,$
2. $\pi(\beta|\mu)$ is $\mathcal{N}(\mu, \lambda_0^{-1}\Sigma^{-1})$, and

3. $\pi(\mu)$ is $\mathcal{N}(\mu_0, \lambda_0^{-1} \Sigma_0^{-1})$.

These assumptions imply that the marginal prior $\pi(\beta)$ is $\mathcal{N}(\mu_0, \lambda_0^{-1} (\Sigma_0^{-1} + \Sigma^{-1}))$.

The posterior distribution of μ is $\mathcal{N}(\mu_1(\phi_0), \lambda_0^{-1} \Sigma_1(\phi_0))$, where

$$\Sigma_1^{-1}(\phi_0) = \Sigma_0 + \left(\Sigma^{-1} + (G^T R^{-1}(\phi_0) G)^{-1} \right)^{-1} \text{ and}$$

$$\mu_1(\phi_0) = \Sigma_1(\phi_0) \left[\left(\Sigma^{-1} + (G^T R^{-1}(\phi_0) G)^{-1} \right)^{-1} \hat{\beta}(\phi_0) + \Sigma_0 \mu_0 \right]$$

for $\hat{\beta}(\phi_0) = (G^T R^{-1}(\phi_0) G)^{-1} G^T R^{-1}(\phi_0) y$.

The posterior distribution of β is $\mathcal{N}(\mu_2(\phi_0), \lambda_0^{-1} \Sigma_2(\phi_0))$, where

$$\Sigma_2^{-1}(\phi_0) = (\Sigma_0^{-1} + \Sigma^{-1})^{-1} + G^T R^{-1}(\phi_0) G \text{ and}$$

$$\mu_2(\phi_0) = \Sigma_2(\phi_0) \left[(G^T R^{-1}(\phi_0) G) \hat{\beta}(\phi_0) + (\Sigma_0^{-1} + \Sigma^{-1})^{-1} \mu_0 \right].$$

A noninformative prior for β , $\pi(\beta) \propto 1$, results in an improper posterior distribution of μ . Note that assuming a noninformative prior for μ , $\pi(\mu) \propto 1$, implies $\pi(\beta) \propto 1$. The posterior distribution of β is as above, with

$$\mu_2(\phi_0) = \hat{\beta}(\phi_0) \text{ and } \Sigma_2(\phi_0) = (G^T R^{-1}(\phi_0) G)^{-1}. \quad (\text{A.1})$$

The predictive distribution of Q future responses $\tilde{Y} = (Y_{\tilde{t}_1}, Y_{\tilde{t}_2}, \dots, Y_{\tilde{t}_Q})^T$ associated with regression matrix $\tilde{G} = [g(\tilde{x}_1) \ g(\tilde{x}_2) \ \dots \ g(\tilde{x}_Q)]^T$ and the Q -variate error vector $\tilde{\varepsilon}_Q = (\varepsilon_{\tilde{t}_1}, \varepsilon_{\tilde{t}_2}, \dots, \varepsilon_{\tilde{t}_Q})^T$ is $\mathcal{N}(\tilde{\mu}(\phi_0), \lambda_0^{-1} \tilde{\Sigma}(\phi_0))$, where

$$\tilde{\mu}(\phi_0) = \tilde{G} \mu_2(\phi_0) + \bar{R}(\phi_0) R^{-1}(\phi_0) (y - G \mu_2(\phi_0)) \text{ and}$$

$$\tilde{\Sigma}(\phi_0) = \tilde{R}(\phi_0) - \bar{R}(\phi_0) R^{-1}(\phi_0) \bar{R}^T(\phi_0) + \tilde{H}(\phi_0) \Sigma_2(\phi_0) \tilde{H}(\phi_0)^T$$

for $\tilde{H}(\phi_0) = \tilde{G} - \bar{R}(\phi_0) R^{-1}(\phi_0) G$. Here, the (i, j) element of $\tilde{R}(\phi_0)$ is given by $r(\tilde{t}_i, \tilde{t}_j | \phi_0)$, and the (i, j) element of $\bar{R}(\phi_0)$ is given by $r(\tilde{t}_i, t_j | \phi_0)$.

(Case 2.) The second verification scenario assumes that ϕ is fixed at ϕ_0 . The Bayesian analysis places a prior distribution on (β, λ) , which is given by the following:

1. $\pi(\beta | \lambda) = \int_{\mathbb{R}^{N_\beta}} \pi(\beta | \mu, \lambda) \pi(\mu | \lambda) d\mu$,
2. $\pi(\beta | \mu, \lambda)$ is $\mathcal{N}(\mu, \lambda^{-1} \Sigma^{-1})$,
3. $\pi(\mu | \lambda)$ is $\mathcal{N}(\mu_0, \lambda^{-1} \Sigma_0^{-1})$, and
4. $\pi(\lambda)$ is $\text{Gamma}(a, b)$.

These assumptions imply that the marginal prior $\pi(\beta | \lambda)$ is $\mathcal{N}(\mu_0, \lambda^{-1} (\Sigma_0^{-1} + \Sigma^{-1}))$.

The posterior distribution of λ is $\text{Gamma}(a_1, b_1(\phi_0))$, where

$$\begin{aligned} a_1 &= (2a + N)/2 \\ b_1(\phi_0) &= \left(2b + \left(y - G\hat{\beta}(\phi_0) \right)^T R^{-1}(\phi_0) \left(y - G\hat{\beta}(\phi_0) \right) \right. \\ &\quad \left. + \left(\hat{\beta}(\phi_0) - \mu_0 \right)^T \Sigma_3^{-1}(\phi_0) \left(\hat{\beta}(\phi_0) - \mu_0 \right) \right) / 2 \end{aligned}$$

for $\Sigma_3(\phi_0) = \Sigma_0^{-1} + \Sigma^{-1} + (G^T R^{-1}(\phi_0) G)^{-1}$.

A noninformative prior for β , $\pi(\beta) \propto 1$, results in the posterior distribution of λ given above, with

$$a_1 = (2a + N - N_\beta)/2 \text{ and } b_1(\phi_0) = \left(2b + \left(y - G\hat{\beta}(\phi_0) \right)^T R^{-1}(\phi_0) \left(y - G\hat{\beta}(\phi_0) \right) \right) / 2. \quad (\text{A.2})$$

A noninformative prior for λ , $\pi(\lambda) \propto (1/\lambda)$, results from taking $a = b = 0$. Note that $\pi(\beta, \lambda) \propto (1/\lambda)$ is the Jeffreys noninformative prior.

We denote the d -variate t distribution having ν degrees of freedom, location vector μ , and scale matrix Σ by $\mathcal{T}_d(\nu, \mu, \Sigma)$. The mean of this distribution is μ if $\nu > 1$ and the covariance matrix of this distribution is $\nu\Sigma/(\nu - 2)$ if $\nu > 2$. The posterior distributions of μ and β are given by

$$\begin{aligned} \pi(\mu|y) &\text{ is } \mathcal{T}_{N_\beta}(2a_1, \mu_1(\phi_0), b_1(\phi_0)\Sigma_1(\phi_0)/a_1) \\ \pi(\beta|y) &\text{ is } \mathcal{T}_{N_\beta}(2a_1, \mu_2(\phi_0), b_1(\phi_0)\Sigma_2(\phi_0)/a_1) . \end{aligned}$$

For the noninformative prior $\pi(\beta) \propto 1$, as before the posterior distribution of μ is improper. The quantities $\mu_2(\phi_0)$, $\Sigma_2(\phi_0)$ from (A.1) and a_1 , $b_1(\phi_0)$ from (A.2) are utilized in the above expression for the posterior distribution of β , where again $a = b = 0$ for $\pi(\lambda) \propto (1/\lambda)$.

The predictive distribution of $\tilde{Y}$ is given by

$$\pi(\tilde{Y}|y) \text{ is } \mathcal{T}_Q \left(2a_1, \tilde{\mu}(\phi_0), b_1(\phi_0)\tilde{\Sigma}(\phi_0)/a_1 \right) .$$

(Case 3.) The third verification scenario allows ϕ to be random. The Bayesian analysis places a prior distribution on (β, λ, ϕ) , which is given by $\pi(\beta, \lambda, \phi) = \pi(\beta, \lambda)\pi(\phi)$ where $\pi(\beta, \lambda)$ is specified as in the previous scenario. The form of $\pi(\phi)$ used in verification testing will be provided in the following section. Our goal in this scenario is to numerically approximate the marginal posterior distributions of λ , μ , and β :

$$\begin{aligned} \pi(\lambda|y) &= \int_{\Phi} \pi(\lambda|y, \phi) \pi(\phi|y) d\phi \\ \pi(\mu|y) &= \int_{\Phi} \pi(\mu|y, \phi) \pi(\phi|y) d\phi \\ \pi(\beta|y) &= \int_{\Phi} \pi(\beta|y, \phi) \pi(\phi|y) d\phi . \end{aligned}$$

The distributions $\pi(\lambda|y, \phi)$, $\pi(\mu|y, \phi)$, and $\pi(\beta|y, \phi)$ are given analytically in the previous scenario. That leaves $\pi(\phi|y)$, which is given as follows:

$$\pi(\phi|y) \propto \frac{\pi(\phi)}{b_1(\phi)^{a_1} \det(R(\phi))^{1/2} \det(G^T R^{-1}(\phi) G)^{1/2} \det(\Sigma_3(\phi))^{1/2}} .$$

Quadrature is used to compute the normalizing constant

$$c(y) = \int_{\Phi} \frac{\pi(\phi) d\phi}{b_1(\phi)^{a_1} \det(R(\phi))^{1/2} \det(G^T R^{-1}(\phi) G)^{1/2} \det(\Sigma_3(\phi))^{1/2}}$$

so that

$$\pi(\phi|y) = \frac{c^{-1}(y) \pi(\phi)}{b_1(\phi)^{a_1} \det(R(\phi))^{1/2} \det(G^T R^{-1}(\phi) G)^{1/2} \det(\Sigma_3(\phi))^{1/2}}.$$

For the noninformative prior $\pi(\beta) \propto 1$,

$$\pi(\phi|y) = \frac{c^{-1}(y) \pi(\phi)}{b_1(\phi)^{a_1} \det(R(\phi))^{1/2} \det(G^T R^{-1}(\phi) G)^{1/2}}$$

for

$$c(y) = \int_{\Phi} \frac{\pi(\phi) d\phi}{b_1(\phi)^{a_1} \det(R(\phi))^{1/2} \det(G^T R^{-1}(\phi) G)^{1/2}},$$

where $a_1, b_1(\phi)$ are taken from (A.2) and $a = b = 0$ for $\pi(\lambda) \propto (1/\lambda)$.

The predictive distribution of $\tilde{Y}$ can also be numerically approximated:

$$\pi(\tilde{Y}|y) = \int_{\Phi} \pi(\tilde{Y}|y, \phi) \pi(\phi|y) d\phi.$$

A.2 Verification Tests

A simple linear regression model, $g(x) = (1, x_1, x_2, \dots, x_M)^T$, is specified for all verification examples. The dimension of β is therefore $N_\beta = M + 1$. A wide range of covariate dimensions are considered, $M \in \{2, 10, 50, 100, 500, 1000\}$.

The verification examples proceed by considering successively more complex specifications of the correlation function $r(t_1, t_2|\phi)$:

1. $r(t_i, t_j) = \delta_{t_i, t_j}$
2. $r(t_i, t_j|\phi) = \delta_{t_i, t_j} + \phi(1 - \delta_{t_i, t_j}), \phi \in [0, 1]$
3. $r(t_i, t_j|\phi) = \phi^{|i-j|}, -1 < \phi < 1$
4. $r(t_i, t_j|\phi) = \exp[-(x_i - x_j)^T D_\phi (x_i - x_j)]$, where $D_\phi = \text{diag}(\phi_1, \dots, \phi_{N_\beta})$ and $\phi_i \geq 0$ for $i = 1, \dots, N_\beta$.

The first case describes the standard regression setting in which errors are uncorrelated. The second case describes equi-correlated errors, for which all pairwise correlations are equal. The third case describes an AR(1) correlation structure for errors associated with observations indexed by time. The pairwise correlation between errors decays as a function of separation in time. The final case specifies a Gaussian correlation for errors associated with observations indexed by covariates. The pairwise correlation between errors decays as a function of weighted distance between the covariate vectors associated with the corresponding observations.

The first three correlation functions admit explicit expressions for the inverse $R^{-1}(\phi)$ and determinant $\det(R(\phi))$,

1. $\mathcal{I}_N$ and 1

2. $\frac{1}{1-\phi} \left[\mathcal{I}_N - \frac{\phi}{1+(N-1)\phi} \mathcal{J}_N \right]$ and $(1-\phi)^{N-1}(1+(N-1)\phi)$

3. $\frac{1}{1-\phi^2} \begin{pmatrix} 1 & -\phi & 0 & 0 & \cdots & 0 & 0 \\ -\phi & 1+\phi^2 & -\phi & 0 & \cdots & 0 & 0 \\ 0 & -\phi & 1+\phi^2 & -\phi & \cdots & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & 0 & \cdots & 1+\phi^2 & -\phi \\ 0 & 0 & 0 & 0 & \cdots & -\phi & 1 \end{pmatrix}$ and $(1-\phi^2)^{N-1}$,

where $\mathcal{I}_N$ is the $N \times N$ identity matrix and $\mathcal{J}_N$ is the $N \times N$ matrix of ones.

For data sample sizes, we consider $\{1, 2, 3, 10, 100, 1000, 10^4, 10^5, 10^6\}$ for the first three correlation functions, and $\{1, 2, 3, 10, 100, 1000\}$ for the fourth correlation function (sample size limited by inverse and determinant calculations). Small sample sizes $\{1, 2, 3\}$ will allow MCMC sampling from heavy-tailed posterior distributions to be compared with known analytical forms (specifically, the multivariate t distributions of the second and third analyses). Large sample sizes allow verification of regression parameter and error variance point estimate convergence ($\hat{\beta} \rightarrow \beta_0$ and $b_1/a_1 \rightarrow \lambda_0^{-1}$) for the first correlation function. Each case will be structured so that the number of random parameters does not exceed the sample size. For example, if $N = 1$, only the first scenario is analyzed with a single random regression parameter β_0 .

Two prior specifications will be considered: (a) $\pi(\mu, \beta, \lambda) \propto (1/\lambda)$ (noninformative) and (b) $\Sigma_0 = q\mathcal{I}_{N_\beta}$ for fixed $0 < q \ll 1$, and $\Sigma = \text{diag}(\{r_i, i = 1, \dots, N_\beta\})$ for fixed $r_i > 0$ (proper). Both cases take $a = b = 0$. For the first correlation function, case (b) allows verification of the convergence of the posterior distribution of μ to the Gaussian distribution having mean vector and covariance matrix

$$\mu_1 = D_\beta \beta_0 + D_\mu \mu_0 \text{ and } \lambda_0^{-1} \Sigma_1 = \text{diag}(\{\lambda_0^{-1}/(q + r_i), i = 1, \dots, N_\beta\})$$

where $D_\beta = \text{diag}(\{r_i/(q + r_i), i = 1, \dots, N_\beta\})$ and $D_\mu = \text{diag}(\{q/(q + r_i), i = 1, \dots, N_\beta\})$. For each of these cases, we take $\pi(\phi) \propto 1$ for the second and third correlation function options. For the fourth option, assume the inputs x satisfy $\ell \leq x \leq u$ and let $\rho_i = \rho_i(\phi) = \exp[-\phi(u_i - \ell_i)^2/4]$ for $\phi \geq 0$. The parameter $\rho_i \in (0, 1]$ represents a correlation length, and we take $\rho_i \sim \text{Beta}(a_\rho, b_\rho)$ with $(a_\rho, b_\rho) = (1, 0.1)$ independently for $i = 1, \dots, M$ as the prior distribution for $\rho = (\rho_1, \dots, \rho_M)$. We then have

$$\pi(\phi) \propto \prod_{i=1}^M [\rho_i(\phi_i)]^{a_\rho} [1 - \rho_i(\phi_i)]^{b_\rho-1} \chi_{[0, \infty)}(\phi_i)$$

where $\chi_{[0, \infty)}(\phi)$ is the characteristic function taking value 1 for $\phi \geq 0$ and 0 for $\phi < 0$.

Bayesian calibration results for these examples are detailed in Section 6.2.6.

Appendix B

Procedure for Running COBRA-TF Studies

This section provides a brief description of how to perform the COBRA-TF studies described in Chapter 7 with the model detailed in Section 2.3 on one of the CASL compute machines using the VERA software environment. This assumes that the steps necessary to log on to one of the various CASL machines, e.g. u233, u235, etc., have been successfully completed.

At present there are two choices for running a VERA VUQ study. One choice is to use the “mainline” codes and utilities that are formally part of the VERA-CS software. This code is tested nightly and adheres to the software quality standards defined by VERA-CS. The alternative is to use an extended capability representing the most current development activities but which is not included in VERA nightly testing. This code and utilities reside in a collection of repositories forked from the mainline. Regardless of which choice is made, the following project directory structure is required:

```
VERA
|
-- Trilinos
| |
| -- packages
|   |
|   -- TriKota
|     |
|     -- Dakota
|
-- COBRA-TF
|
-- VERAINExt
|
-- LIMExt
|
-- PSSDriversExt
|
-- VUQDemos
```

Table B.1 shows the locations from which to clone each repository for either choice.

Table B.1: Code repository locations for either “mainline” or “vuq_fork” versions of CASL VUQ.

Repository Name	mainline	vuq_fork
VERA	casl-dev:/git-root/VERA	casl-dev:/git-root/VERA
Trilinos	casl-dev:/git-root/Trilinos	casl-dev:/git-root/Trilinos
Dakota	casl-dev:/git-root/Dakota	casl-dev:/git-root/Dakota
COBRA-TF	casl-dev:/git-root/COBRA-TF	casl-dev:/home/vuq-root/COBRA-TF_vuq_fork
VERAInExt	casl-dev:/git-root/VERAInExt	casl-dev:/git-root/VERAInExt
LIMEExt	casl-dev:/git-root/LIMEExt	casl-dev:/git-root/LIMEExt
PSSDriversExt	casl-dev:/git-root/casl_vripss	casl-dev:/home/vuq-root/casl_vripss_vuq_fork
VUQDemos	casl-dev:/git-root/VUQDemos	casl-dev:/home/vuq-root/VUQDemos_vuq_fork

For example, Listing B.1 shows the steps to set up a CASL VUQ project on a CASL machine using the “vuq_fork” collection of codes and tools. After the project has been created, it can be built in a variety of ways. Probably the easiest approach and the one guaranteed to work is to make use of the VERA checkin-test utility script. This should be done within a build directory separate from the project source directory created above. For example, to build an optimized version of the project that runs in serial the steps shown in Listing B.2 would be followed. This will perform a sequence of steps which include configuring (i.e. creating Makefiles), building (i.e. invoking make) and testing the project (i.e. invoking ctest). The option, “test-categories=WEEKLY”, enables the most comprehensive testing and can include some tests that require on the order of an hour to complete. A list of tests can be obtained by going into the SERIAL_RELEASE subdirectory and invoking “ctest -N”, and a particular test can be run by specifying its name or regular expression encompassing its name, e.g. “ctest -R single_assembly”. If the project source directory is created in a location other than the top-level user home directory, i.e. “~/VERA”, then its location must be specified in the invocation of the checkin-test script by adding the option, “-src-dir=\$PATH_TO_TOP_LEVEL_VERA_DIR”, and the path must be absolute, e.g. “home/\$USER_ID/Projects/CASL/VERA”. Finally, it should be noted that the checkin-test utility script currently requires that it be invoked from the second level of a build directory tree. Any build directory structure of the following form would work, “~/SOME_LEVEL_1_DIR/LEVEL_2_DIR” with the script residing and invoked from the LEVEL_2_DIR.

Listing B.1: Steps to setup a vuq_fork CASL VUQ project on one of the CASL machines

```
2 # Checkout top-level VERA source directory
  git clone casl-dev:/git-root/VERA
4
  # Checkout Trilinos under VERA
6 cd VERA
  git clone casl-dev:/git-root/Trilinos
8
  # Checkout Dakota in proper subdirectory under Trilinos
10 cd Trilinos/packages/TriKota
  git clone casl-dev:/git-root/Dakota
12
  # Checkout Cobra-TF Thermal Hydraulics code under VERA
14 cd ../../..
  git clone casl-dev:/home/vuq-root/COBRA_TF_vuq_fork COBRA-TF
16
  # Checkout VERA Common Input toolset under VERA
18 git clone casl-dev:/git-root/VERAInExt
20
  # Checkout LIME multi-physics coupling driver under VERA
  git clone casl-dev:/git-root/LIMEExt
22
  # Checkout VERA interfaces to CASL application codes under VERA
24 git clone casl-dev:/home/vuq-root/casl_vripss_vuq_fork PSSDriversExt
26
  # Checkout CASL VUQ toolset under VERA
  git clone casl-dev:/home/vuq-root/VUQDemos_vuq_fork VUQDemos
```

Listing B.2: Steps to build a vuq_fork CASL VUQ project on one of the CASL machines

```
1
  # Create a build directory
3 mkdir -p CASL_BUILDS/VUQ_BUILD
  cd CASL_BUILDS/VUQ_BUILD
5
  # Create a link to the checkin-test utility script in the source project directory
7 ln -s ~/VERA/cmake/ctest/drivers/fissile4/checkin-test-vera.sh
9
  # Invoke the script to build an optimized serial version of the project
  ./checkin-test-vera.sh --extra-repos-type=Nightly --test-categories=WEEKLY \
11 --default-builds=SERIAL_RELEASE --enable-packages=VUQDemos --local-do-all
```

Bibliography

- [1] 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, and D. M. Vigil. Dakota, a multilevel parallel object-oriented framework for design optimization, parameter estimation, uncertainty quantification, and sensitivity analysis: Version 5.4 user's manual. Technical Report SAND2010-2183, Sandia National Laboratories, Albuquerque, NM, Updated Nov. 2013. Available online from <http://dakota.sandia.gov/documentation.html>.
- [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, and D. M. Vigil. Dakota, a multilevel parallel object-oriented framework for design optimization, parameter estimation, uncertainty quantification, and sensitivity analysis: Version 5.4 reference manual. Technical Report SAND2010-2184, Sandia National Laboratories, Albuquerque, NM, Updated Nov. 2013. Available online from <http://dakota.sandia.gov/documentation.html>.
- [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, and D. M. Vigil. Dakota, a multilevel parallel object-oriented framework for design optimization, parameter estimation, uncertainty quantification, and sensitivity analysis: Version 5.4 theory manual. Technical Report SAND2011-9106, Sandia National Laboratories, Albuquerque, NM, Updated Nov. 2013. Available online from <http://dakota.sandia.gov/documentation.html>.
- [4] J. S. Arora. *Introduction to Optimum Design*. McGraw-Hill, New York, 1989.
- [5] M. N. Avramova. CTF: A thermal hydraulic sub-channel code for LWR transient analyses, users manual. Technical report, Pennsylvania State University, February 2009.
- [6] L. S. Bastos and A. O'Hagan. Diagnostics for Gaussian process emulators. *Technometrics*, 51:425–438, 2009.
- [7] G. E. P. Box and D. W. Behnken. Some new three level designs for the study of quantitative variables. *Technometrics*, 2:455–475, 1958.
- [8] G. E. P. Box and K. B. Wilson. On the experimental attainment of optimum conditions (with discussion). *Journal of the Royal Statistical Society Series B*, 13:1–45, 1951.
- [9] N. Cressie. *Statistics of Spatial Data*. John Wiley and Sons, New York, 1991.
- [10] K. R. Dalbey, A. A. Giunta, M. D. Richards, E. C. Cyr, L. P. Swiler, S. L. Brown, M. S. Eldred, and B. M. Adams. Surfpack user's manual: Version 1.1. Technical report, Sandia National

- Laboratories, Albuquerque, NM, 2006. Available online from <http://dakota.sandia.gov/packages/Surfpack>.
- [11] A. Gelman, J. B. Carlin, H. S. Stern, and D. B. Rubin. *Bayesian Data Analysis*. Chapman & Hall/CRC, Boca Raton, 1997.
 - [12] R. Ghanem and J. R. Red-Horse. Propagation of probabilistic uncertainty in complex physical systems using a stochastic finite element technique. *Physica D*, 133:137–144, 1999.
 - [13] R. G. Ghanem and P. D. Spanos. *Stochastic Finite Elements: A Spectral Approach*. Springer-Verlag, New York, 1991.
 - [14] P. E. Gill, W. Murray, and M. H. Wright. *Practical Optimization*. Academic Press, San Diego, CA, 1981.
 - [15] Heikki Haario, Marko Laine, Antonietta Mira, and Eero Saksman. DRAM: Efficient adaptive MCMC. *Statistics and Computing*, 16:339–354, 2006.
 - [16] R. T. Haftka and Z. Gurdal. *Elements of Structural Optimization*. Kluwer, Boston, 1992.
 - [17] W. E. Hart. The Coliny project. Web site, 2007.
 - [18] A. S. Hedayat, N. J. A. Sloane, and J. Stufken. *Orthogonal Arrays: Theory and Applications*. Springer, New York, 1999.
 - [19] D. Higdon, J. Gattiker, B. Williams, and M. Rightley. Computer model calibration using high-dimensional output. *Journal of the American Statistical Association*, 103(482):570–583, 2008.
 - [20] E. T. Jaynes and G. Larry Bretthorst. *Probability theory : the logic of science*. Cambridge University Press, Cambridge, UK; New York, NY, 2003.
 - [21] C. T. Kelley. *Implicit Filtering*. SIAM, 2011.
 - [22] M. C. Kennedy and A. O’Hagan. Bayesian calibration of computer models. *Journal of the Royal Statistical Society*, 63:425–464, 2001.
 - [23] J. L. Loeppky, J. Sacks, and W. J. Welch. Choosing the sample size of a computer experiment: A practical guide. *Technometrics*, 51:366–376, 2009.
 - [24] 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.
 - [25] M. D. Morris. Factorial sampling plans for preliminary computational experiments. *Technometrics*, 33(2):161–174, 1991.
 - [26] V. A. Mousseau et al. VUQ strategy. Technical Report CASL-U-2014-XXXX-YYY, Oak Ridge National Laboratory, March 2014. In preparation.
 - [27] R. H. Myers and D. C. Montgomery. *Response Surface Methodology: Process and Product Optimization Using Designed Experiments*. John Wiley & Sons, Inc., New York, 1995.

- [28] J. Nocedal and Wright S. J. *Numerical Optimization*. Springer Series in Operations Research. Springer, New York, 1999.
- [29] W.T. Nutt and G.B. Wallis. Evaluation of nuclear safety from the outputs of computer codes in the presence of uncertainties. *Reliability Engineering and System Safety*, 83:57–77, 2004.
- [30] W. L. Oberkampf, M. M. Pilch, and T. G. Trucano. Predictive capability maturity model for computational modeling and simulation. Technical Report SAND2007-5948, Sandia National Laboratories, October 2007.
- [31] Scott Palmtag. Coupled single assembly solution with VERA (problem 6). Technical Report CASL-U-2013-0150-000, Oak Ridge National Laboratory, 2013.
- [32] Robert Salko. CTF list of global variables. Github code repository, 2013.
- [33] A. Saltelli, S. Tarantola, F. Campolongo, and M. Ratto. *Sensitivity Analysis in Practice: A Guide to Assessing Scientific Models*. John Wiley & Sons, 2004.
- [34] R. C. Smith. *Uncertainty Quantification: Theory, Implementation and Applications*. SIAM, Philadelphia, PA, 2014.
- [35] A. Solonen, P. Ollinaho, M. Laine, H. Haario, J. Tamminen, and H. Järvinen. Efficient MCMC for climate model parameter estimation: parallel adaptive chains and early rejection. *Bayesian Analysis*, 7(3):715–736, 2012.
- [36] R. Sues, M. Aminpour, and Y. Shin. Reliability-based multidisciplinary optimization for aerospace systems. In *Proc. 42nd AIAA/ASME/ASCE/AHS/ASC Structures, Structural Dynamics, and Materials Conference*, number AIAA-2001-1521, Seattle, WA, April 16-19 2001.
- [37] L.P. Swiler and V.J. Romero. JANNAP V&V guide: A survey of advanced probabilistic uncertainty propagation and sensitivity analysis methods. Technical Report SAND2013-5456P, Sandia National Laboratories, Albuquerque, NM, 2013.
- [38] G. N. Vanderplaats. *Numerical Optimization Techniques for Engineering Design: With Applications*. McGraw-Hill, New York, 1984.
- [39] J.A. Vrugt, C.J.F. Ter Braak, M. P. Clark, J. M. Hyman, and B. A. Robinson. Treatment of input uncertainty in hydrological modeling: Doing hydrology backward with markov chain monte carlo, simulation. *Water Research Resources*, 44(12):W00B09, doi:10.1029/2007WR006720, 2008.
- [40] J.A. Vrugt, C.J.F. Ter Braak, C. G. H. Diks, B. A. Robinson, J. M. Hyman, and D. Higdon. Accelerating markov chain monte carlo simulation by differential evolution with self-adaptive randomized subspace sampling. *International Journal of Nonlinear Sciences and Numerical Simulation*, 10(3):273–290, 2009.
- [41] S.S Wilks. Determination of sample sizes for setting tolerance limits. *The Annals of Mathematical Statistics*, 12(1):91–96, 1941.

- [42] 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.

DISTRIBUTION:

1 MS 0899 Technical Library, 9536 (electronic copy)

