

**Report to the National
Measurement System
Policy Unit, Department of
Trade and Industry**

**Uncertainty Evaluation in
Continuous Modelling**

G J Lord and L Wright

December 2003

Uncertainty Evaluation in Continuous Modelling

G J Lord* and L Wright†

*Department of Mathematics, Heriot-Watt University

†Centre for Mathematics and Scientific Computing, NPL

December 2003

ABSTRACT

This report describes methods for evaluating the uncertainties associated with the results of continuous models, and provides case studies of the application of these methods to real problems. The methods include sampling methods such as Monte Carlo simulation, analytical methods such as stochastic differential equations, and methods used for particular types of problem such as Structural Energy Analysis. The case studies apply some of the methods to three real metrology problems and demonstrate the benefits and disadvantages of the methods.

© Crown Copyright 2003
Reproduced by Permission of the Controller of HMSO

ISSN 1471-0005

National Physical Laboratory
Queens Road, Teddington, Middlesex, TW11 0LW

Extracts from this report may be reproduced provided the source is acknowledged and
the extract is not taken out of context

Approved on behalf of the Managing Director, NPL
by Dave Rayner, Head of the Centre for Mathematics and Scientific Computing

Contents

1	Introduction	1
1.1	Existing application and research areas	1
1.2	Structure of the report	3
2	Issues affecting uncertainty evaluation in continuous models	4
2.1	The GUM approach.....	4
2.2	A supplement to the GUM	6
2.3	Sources of uncertainty.....	7
2.4	Correlation.....	8
2.5	Sensitivity.....	8
2.6	Validation	9
2.7	Definitions of probability and statistics terms	10
3	Methods of uncertainty evaluation.....	13
4	Sampling methods.....	14
4.1	Design of experiments and “extreme value” methods	14
4.2	Monte Carlo simulation.....	17
4.2.1	Constructing the distribution function	18
4.2.2	Choosing M	19
4.2.3	An example of MCS.....	20
4.3	Constrained Monte Carlo simulation	24
4.3.1	Stratified sampling	24
4.3.2	Latin Hypercube sampling	26
4.3.3	An example of constrained MCS methods.....	29
5	Analytical methods.....	32
5.1	Stochastic differential equations	32
5.1.1	Calculating statistics and estimating parameters using SDEs.....	34
5.2	Emulators and response surface methodologies	36
5.3	Series expansion methods	37
6	Methods developed for particular applications	40
6.1	Statistical Energy Analysis	40
6.2	Most Probable Point methods	40
7	Case studies.....	42
7.1	Case Study 1: Strain measurement in blades	43
7.1.1	The physical problem.....	43
7.1.2	The mathematical model	43
7.1.3	Sources of uncertainty.....	43
7.1.4	Initial study.....	44
7.1.5	Choice of calculation method.....	45
7.1.6	Results of the calculations.....	45
7.1.7	Combining the uncertainties from different sources	48
7.1.8	Conclusions	49
7.2	Case Study 2: Determination of the time constant of a calorimeter	50
7.2.1	The physical problem.....	50
7.2.2	The mathematical model and choice of calculation method	50
7.2.3	Sources of uncertainty.....	51
7.2.4	Results of the calculations.....	52
7.2.5	Conclusions and extensions	52
7.3	Case Study 3: Laser Flash Thermal Diffusivity.....	54

7.3.1	The physical problem	54
7.3.2	The mathematical model	54
7.3.3	Sources of uncertainty	55
7.3.4	Choice of calculation method.....	55
7.3.5	Results of the calculations.....	56
7.3.6	Conclusions and extensions	58
8	Conclusions	59
9	References	60
10	Appendix I: Details of stochastic differential equation techniques.....	63
10.1.1	Brownian Motion	63
10.1.2	Stochastic calculus	64
10.1.3	Useful formulae and extensions	66
10.1.4	Numerics for SDEs.....	69
10.1.5	Convergence and stability	70
10.1.6	Further reading	73

1 Introduction

This report describes methods for evaluating the uncertainties associated with the results of continuous models, and provides case studies of the application of these methods to real problems. It also examines reasons why uncertainties are less commonly provided for the results of continuous modelling, and highlights the difficulties of applying to continuous modelling problems some of the approaches to uncertainty evaluation commonly used in discrete modelling.

The contents of this report were partly motivated by the results of a survey of SSfM club members. The survey requested details of respondents' continuous modelling work, including the reasons why they commonly undertook such work. The most popular responses were "for processing measurement data", "for experimental or equipment design", and "for a deeper understanding of the measurement process".

The aim of processing measurement data is generally to estimate the value of a quantity that is to be reported as a measurement result. Since the quantity depends (through a model) on other quantities, the values of which are only inexactly known, the measurement result will itself be inexact. For its interpretation and subsequent use, the result is therefore reported with an associated uncertainty that quantifies the inexactness. The model is the basis of the evaluation of the uncertainty associated with the measurement result.

The purpose of experimental design is usually to investigate the possible sources of uncertainty that may affect a measurement result. This investigation needs to be able to produce an uncertainty associated with the measurement result so that different experimental designs can be compared, and so uncertainty evaluation is a key part of this application of continuous modelling.

The use of continuous models as investigative tools is common when there are phenomena that have well-understood models but that are not necessarily easy to measure directly in the system being simulated. The models can give insights into how the different phenomena affect the measurement result. The simulation of these phenomena means that if the model is used to evaluate uncertainties, it would be possible to evaluate the contributions to the uncertainty from the unmeasured effects, leading to an improved overall uncertainty value.

These cases show that uncertainty evaluation is crucial to many of the applications for continuous models that are commonly used in metrology. This report investigates factors that need consideration when using continuous models for uncertainty evaluation, describes methods that are suitable for such problems, and illustrates their use with three case studies from various areas of metrology.

This report is strongly linked to the report "Model Validation In Continuous Modelling" [1]. In that report, parameter uncertainty is briefly considered as one of the factors that contribute to the inexactness of the results of continuous models, but the issues surrounding uncertainty evaluation are not examined in detail.

1.1 Existing application and research areas

A significant amount of work on uncertainties and sensitivity analysis in continuous modelling has been done in the field of structural mechanics [2-10]. The main impetus for this is prediction of reliability. Many structural mechanics models are used to design

components, objects and large-scale structures, e.g. bridges, aeroplanes, and buildings, for which safety and reliability are prime concerns. The safety requirements are usually that 99.9% of the possible values of some quantity (stress, displacement, temperature etc.) lie within a prescribed range. To show that the requirement is met, it is necessary to consider the tails of the distribution for the value of the quantity.

Since many of these models are of large structures and therefore require significant computation to run, much work has gone into identifying suitable methods for guaranteeing reliability using the minimum amount of processor time. Clearly the fewer times a model needs to be run for the construction of a distribution function, the less precise and accurate the results will be, but often the nature of the question of interest makes this trade-off acceptable.

This focus of application of the research has led to techniques being developed that are only applicable to structural dynamics problems, for instance structural energy analysis techniques. These methods are explained in section 6. Whilst the methods were developed for structural problems, they are likely to be applicable to any problem requiring the determination of the fundamental frequencies of an object that is composed of multiple parts.

As well as application-specific methods, more general techniques have been applied successfully to continuous modelling problems. These include Monte Carlo simulation (MCS) [11-13] (see section 4.2), other sampling methods (see section 4.3) such as Latin Hypercube sampling [13-15], Stratified sampling [13], and Adaptive Importance sampling [9]. A lot of the work [14,16] has been to compare these different sampling methods in order to explore the most efficient way of obtaining reliable information about the modelled situation.

Sampling techniques treat the model as a black box and ignore potentially useful features of the model such as linearity, symmetry, or knowledge of derivatives. Exploitation of these features often requires much more challenging mathematics than sampling methods, but some techniques have been applied to continuous models. Stochastic differential equations are used extensively in financial applications, and some researchers [17] have applied techniques from stochastic differential equation theory to problems involving uncertainty analysis. These techniques are quite demanding to understand and implement, but they are still useful in some cases. They are explained further in section 5.1. Other work [18] has used series expansion methods and stochastic finite element analysis formulation to solve problems. These methods are explained in section 5.3.

Some researchers [3-8, 10] have created interfaces for common finite element packages that use various techniques to produce coverage intervals for the results of interest. Examples of such packages include SLang [3, 4], DAKOTA [5, 6], NESSUS [7, 8], and ProFES [10]. The most commonly used sampling method is Monte Carlo simulation [2], but other methods are applied. NESSUS uses most probable point methods (see section 6.2) to develop a density function for a given quantity of interest. DAKOTA supports Latin Hypercube sampling and Monte Carlo simulation for uncertainty analysis, as well as various reliability techniques (including most probable point methods), and robustness techniques that provide coverage intervals for a quantity when little is known about the uncertainties of the input quantities. SLang is a tool for carrying out Monte Carlo simulations with finite element packages and can accept a wide range of quantities as uncertain parameters.

Some methods are not covered in this report because they are not sufficiently developed for general applicability. In particular, innovative methods such as fuzzy finite elements are being investigated as tools for uncertainty evaluation, but the research is still quite a long way from practical application.

1.2 Structure of the report

This report is split into eight main sections. Section 2 outlines some of the main issues affecting uncertainty evaluation in continuous models, and explains why the procedure predominantly endorsed in the GUM [19] is not appropriate for the majority of continuous models. The section also includes definitions of some probabilistic and statistical terms that will be used elsewhere in the report.

Section 3 explains the main classes of method that exist and looks at issues affecting choice of method. Sections 4, 5 and 6 discuss particular classes of method, using worked examples to demonstrate and compare techniques. These sections also examine the pros and cons of the different methods and point out features that need to be taken into account when applying the methods. The classes examined are sampling methods, analytical methods, and methods developed for particular applications. The mathematics describing one of the analytical methods, stochastic differential equations, is particularly involved, so it is outlined in the appendix, section 10.

Section 7 consists of three case studies demonstrating the methods. The case studies are drawn from three different areas of metrology, and use different techniques to evaluate the relevant uncertainties. Each case study applies at least one of the methods outlined, and where possible the studies compare several different methods.

2 Issues affecting uncertainty evaluation in continuous models

This section gives an insight into some of the issues affecting uncertainty evaluation using continuous models. It describes some of the features of real problems that a simulation method needs to preserve, such as correlation. It also includes a brief introduction to some probabilistic and statistical concepts that will be used elsewhere.

2.1 The GUM approach

The most common method of calculating uncertainty in metrology is probably that outlined in Clause 8 of the GUM [19]. Once the model linking the input quantities and output quantities has been identified, and probability density functions (pdfs) have been assigned to the input quantities, the main steps in the procedure are:

1. Get the means, standard uncertainties, and if necessary covariances from the pdfs for the values of the input quantities, or joint pdfs for mutually dependent input quantities.
2. Form the partial derivatives of the model with respect to the input quantities.
3. Evaluate the model at the estimates of the input quantities to give the measurement result.
4. Evaluate the partial derivatives at the estimates of the input quantities to give the sensitivity coefficients.
5. Combine the values of the standard uncertainties, the covariances, and the model sensitivity coefficients according to GUM Formula (10) to give the combined standard uncertainty associated with the measurement result.
6. Calculate the effective degrees of freedom ν using the Welch-Satterthwaite formula (GUM Equation (G.2b))
7. Compute the expanded uncertainty and a coverage interval for the value of the output quantity by forming the appropriate multiple of the combined standard uncertainty depending on the value of ν . If $\nu = \infty$, a Gaussian distribution is used to assign the multiple, and otherwise a t -distribution with ν degrees of freedom is used.

This process is based on applying the law of propagation of uncertainty and invoking the Central Limit Theorem to assign a distribution for the value of the output quantity. The procedure relies on a number of assumptions:

- The linearisation of the model is adequate.
- The representation of the distribution for the value of the output quantity in terms of a Gaussian distribution or a t -distribution is adequate.
- The Welch-Satterthwaite formula is adequate for the problem.
- The input quantities are uncorrelated if ν is finite.

A feature of these assumptions is that they are “uncontrolled”: it is difficult to quantify them or put them right. Other methods, such as MCS, also make assumptions but these are “controllable” (e.g. by increasing the number of MC trials).

A number of common features of continuous models make the above GUM methodology unsuitable. The main features are:

- It is often very difficult to obtain the partial derivatives required in step 2 for a continuous model. Often the boundary and initial conditions are the input quantities, and the links between the output quantity and these input quantities are not expressed in simple closed form. Obtaining the partial derivatives by numerical methods, such as finite difference methods, is often time-consuming, particularly for problems with a large number of input quantities.
- Frequently the output quantity for a continuous model is the location and value of the maximum of some quantity, for instance the peak stress or the maximum temperature. This type of result can produce an expression for the output quantity that is not differentiable, and so no partial derivatives can be obtained in step 2.
- Some continuous models do not depend on their input quantities in a straightforward manner [1, section 4.2], for instance multiple solutions may exist if the model exhibits hysteresis. This means that the linearity requirement for the model is unlikely to hold, and the Central Limit Theorem is unlikely to apply.

These points are not meant to imply that the GUM methodology is unsuitable for all continuous models. Many models depend on their input quantities linearly and satisfy the other conditions necessary for the GUM methodology to apply. For example, consider one-dimensional heat transfer along a bar of unit length and negligible cross-section. Suppose that the one end of the bar is held at an inexact fixed temperature and that the other end is losing heat at an inexact rate. The temperature profile will satisfy

$$\frac{\partial}{\partial x} \left(\lambda \frac{\partial T}{\partial x} \right) = 0, \quad 0 \leq x \leq 1,$$

$$T(0) = T_0, \quad \lambda \frac{\partial T}{\partial x} \bigg|_{x=1} = Q_1,$$

where T_0 and Q_1 are independent inexact quantities and λ is the (exactly known) thermal conductivity. The analytical solution to this problem is $T = T_0 + x(Q_1/\lambda)$. Suppose that T_0 and Q_1 have Gaussian distributions, with means 0 °C and 1 Wm⁻² respectively, and standard deviations 0.2 °C and 0.1 Wm⁻² respectively. If the output quantity T^* is the temperature at the midpoint of the bar, then $T^* = T_0 + Q_1/(2\lambda)$ and the GUM methodology can be applied to this problem. The steps in the methodology give:

1. $T_0 \sim N(0, 0.04)$ and $Q_1 \sim N(1, 0.01)$
2. $\frac{\partial T^*}{\partial T_0} = 1, \quad \frac{\partial T^*}{\partial Q_1} = \frac{1}{2\lambda}$
3. Measurement result is $\bar{T}^* = 1/2\lambda$
4. Sensitivity coefficients are 1 for T_0 and $1/(2\lambda)$ for Q_1 .
5. Combined standard uncertainty associated with the measurement result is $u_c^2 = 1 \times 0.04 + 0.01/(4\lambda^2)$
6. The effective degrees of freedom ν is taken to be infinite since it is assumed that

the distributions of T_0 and Q_1 are known exactly.

7. The expanded uncertainty for a 95% coverage interval for the value of T^* is $U_c = 1.96 \times [0.04 + 0.0025/\lambda^2]^{1/2}$, a coverage interval for the value of the output quantity is $[1/(2\lambda) - U_c, 1/(2\lambda) + U_c]$, and T^* will be taken to have a Gaussian distribution.

It can be shown that if X_1 and X_2 are independent normally distributed variables with means μ_1 and μ_2 respectively and standard deviations σ_1 and σ_2 respectively, and if a is some constant, then $X_1 + aX_2$ will be normally distributed with mean $\mu_1 + a\mu_2$ and standard deviation $(\sigma_1^2 + a^2\sigma_2^2)^{1/2}$. This result agrees with the values given in steps 3, 5, and 7 above.

The GUM methodology should always be considered as an initial approach as it is straightforward and well-established. In some cases it may be possible to produce a simplified version of a continuous model that can have the GUM methodology applied to it, to give a rough estimate of the uncertainty before proceeding with a more complex version of the model. The methodology will not be discussed further in this report since it is described extensively in the GUM [19].

2.2 A supplement to the GUM

Recently, a Supplemental Guide to the GUM [11] has recognised that for an increasing number of metrology problems the assumptions necessary for the application of the method outlined above do not apply. As a result, efforts have been made to introduce a more widely applicable methodology for evaluating uncertainties. This methodology can be used for validating uncertainties evaluated using the GUM approach as well as for evaluating them in cases where the approach does not apply.

The methodology is based on the concept of propagation of distributions. The methodology does not rely on any of the assumptions mentioned in section 2.1, and for cases that fulfil the assumptions outlined in section 2.1 the two methodologies will produce comparable results.

The supplement to the GUM restricts its attention to models with a single output quantity rather than vector or field quantities, but another Supplemental Guide covering models with more than one output quantity is currently in preparation. The method outlined for propagation of distributions relies on Monte Carlo simulation. MCS is described more fully in section 4.2 and so will not be described here, but MCS produces an approximation to the pdf of the output quantity from which a mean, standard uncertainty, and coverage intervals can be obtained. A variation of this approach can be used to obtain something similar to sensitivity coefficients.

Whilst this approach is sufficiently general to make it applicable to many continuous models, it still has a drawback. Many continuous models require the solution of large systems of simultaneous equations during their solution procedure. Often these equations are solved using iterative methods, and in some cases it may be necessary to solve a large number of simultaneous equations at a number of time steps. This solution procedure can lead to a lengthy run time for some continuous models, since the solution of a single iteration for even a medium-sized model can easily require millions of floating point operations per time step. This lengthy run time can make methods such as MCS, which require repeated model runs, prohibitively expensive.

Additionally, many continuous model results are correlated multi-variate quantities,

since they are samples of an approximation to a continuous variable. The method as outlined in the supplement to the GUM is not directly applicable in such cases.

2.3 Sources of uncertainty

In the report on continuous model validation [1], five main sources of error and uncertainty were identified. They were:

- Modelling error: the failure of the chosen continuous equation (generally either a differential or integral equation) and boundary or initial conditions to describe reality.
- Space discretization error: error generated by solving a continuous problem on a discrete mesh.
- Time discretization error: error generated (and accumulated) by solving a continuous problem at discrete time steps.
- Parameter errors: errors generated by inaccurate input quantities.
- Linear algebra errors: errors generated during the solution of the discretized system.

These sources are illustrated in relation to the modelling process in figure 1. The steps in the modelling process are listed on the left, and the sources of uncertainty affecting each step are listed on the right.

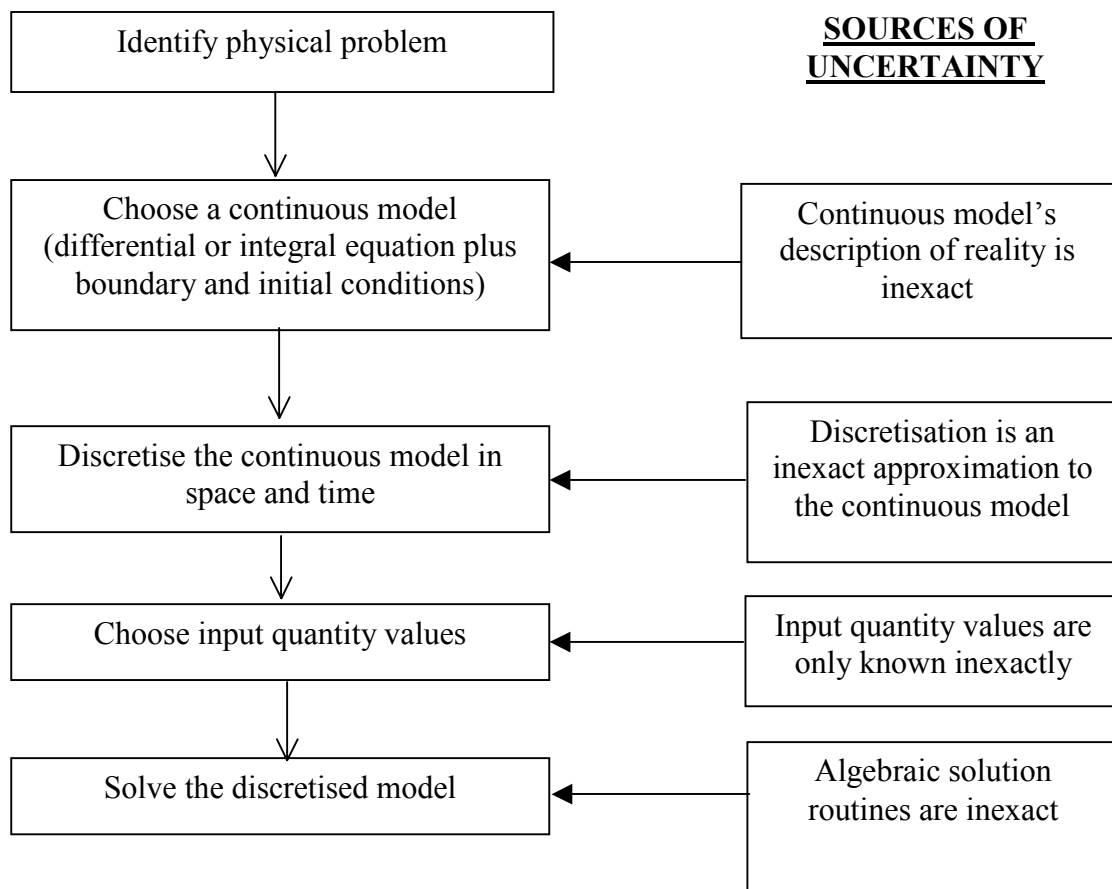


Figure 1: An illustration of the sources of uncertainty and how they relate to the modelling process.

All of these sources will contribute to the overall difference between the model and reality. The uncertainty associated with the model's output quantity is usually regarded as being caused by the inexactness of the input parameters and solution procedures of the continuous model, so that the model's inexact description of reality (which is usually unquantifiable) can be neglected.

Model validation techniques [1] can be used to estimate and reduce the discretisation errors caused by approximating a continuous quantity at discrete points. Software testing methodologies [20, 21] can be used to identify inexact algebraic routines. It is assumed that a model in use for uncertainty evaluation has been validated and tested to a sufficient degree that these sources of uncertainty are negligible when compared to the uncertainty caused by the use of inexact input quantities.

This report describes methods for calculating the contribution to the uncertainty from the inexactness of the input quantities. It should be remembered that input parameters are not the only cause of the uncertainty associated with the output quantity, and although good model validation and software testing procedures will minimise the other sources of uncertainty, they still need to be considered.

2.4 Correlation

Correlation occurs when one or more quantities are linked in some way: circumstances that produce a change in one quantity produce a corresponding change in the other. For instance, most physical properties of materials are temperature-dependent to some degree, so different model parameters will be linked via the ambient temperature. This linking means that the parameters are not independent, and their correlation will need to be taken into account when determining the output uncertainty.

Correlated input quantities are covered in section 5.2 of the GUM by including terms involving the correlation coefficients into the equation for combined uncertainty. Any other method used for uncertainty evaluation needs to be able to account for correlation as it can be an important feature of the physical situation. For instance, the sampling methods that are described in section 4 account for correlation between the input quantities by sampling from their joint pdf.

Correlation can also occur between output quantities. It is almost certain that multiple results from a single continuous model will be correlated to some extent. This correlation occurs because in general the solution to a partial differential equation at any point is dependent on all the input quantities and conditions, and the continuous model results that approximate this solution share this dependence. The correlation means that if multiple results are of interest, a covariance matrix or a joint pdf should be supplied when reporting the uncertainty.

2.5 Sensitivity

The partial derivatives of the model with respect to the input quantities evaluated at the estimates of those quantities, as determined in step 4 of the process described in section 2.1 above, are called sensitivity coefficients. These are of use not only as a step towards the evaluation of a combined uncertainty, but they can give the metrologist an indication of which input uncertainties it is worth trying to reduce. Reduction of a significant input uncertainty with a large sensitivity coefficient can produce a useful reduction in the combined uncertainty of the output. This means that it can be worthwhile to evaluate the sensitivity coefficients even when they are not explicitly

required for uncertainty evaluation.

The evaluation of these coefficients can be particularly problematic for some continuous models. As has been shown in the continuous model validation report [3], dependence of model results on model parameters can take some unusual forms. Multiple solutions can exist, for example due to hysteresis, which is clearly problematic.

Some methods of uncertainty evaluation aim to provide sensitivity coefficients as well. This is not a subject that will be explored in depth in this report, but sensitivity analysis is a subject [22] that is growing in popularity as increased computational power makes study of complicated models more tractable.

A fairly recently-developed technique that may be of use in this context is automatic differentiation [23-25]. Automatic differentiation uses the chain rule of differentiation and simple rules for differentiation of key functions such as polynomials and trigonometric functions to construct the partial derivatives of more complicated compound functions with respect to their input quantities. The technique uses the modelling software source code as a definition of the model in terms of these key functions and calculates derivative values from repeated application of the chain rule. It does not calculate analytical formulae for derivatives, but instead it calculates derivative values for given input values. The simplicity of the basic rules and the universality of their application mean that the technique is flexible, thorough, and can be applied to a very wide range of problems.

The technique works on single-valued and vector-valued quantities and can be used to determine higher order derivatives if required. At present it is used for calculating derivatives for optimisation problems and Newton-Raphson type solvers, but its use as a tool for calculating sensitivity coefficients is increasing [23, 24]. The main drawback of the technique for some applications in continuous modelling is its requirement of source code access, so at present it cannot be used in conjunction with proprietary packages.

2.6 Validation

The SSfM Best Practice Guide (BPG) on Discrete Model Validation [26] has advice on validation of the statistical models used to characterise the uncertainty of the input quantities (in the discrete modelling case, the input quantities are assumed to be measurement data). The advice given is relevant to continuous models as well, since there is usually no difference between the nature of the input quantities for the two types of model.

The advice consists of a number of points that need to be taken into consideration when choosing statistical models for the input quantities. These points are:

- Check that any quantities being taken to be exact values have uncertainties that can be regarded as insignificant,
- Check that any uncertainties that are dependent on the magnitude of the input quantity (for instance relative errors) are treated as such,
- Consider the independence of the input quantities.
- Consider the likely behaviour of the input quantity when assigning a pdf, particularly its symmetry, peaks, and behaviour at the tails of the distribution.

This is particularly important for quantities that have known limits, for instance many physical quantities have to be positive and so any distribution with a non-zero probability for negative values is unsuitable.

The BPG also suggests some techniques for validating the statistical model, such as:

- Review by an expert in the relevant area of metrology.
- Review by an expert modeller.
- Use of numerical simulations to check the effects of the assigned distributions on the results.

The BPG also notes that it can be difficult to detect a poor statistical model. A poor model of the physical problem can be identified by its poor reproduction of expected behaviour (measured or calculated), but a poor statistical model is less likely to be identified unless a large collection of model results and measured data are available.

2.7 Definitions of probability and statistics terms

This section contains definitions and explanations of terms that are used throughout this report. It can be skipped by readers with a good knowledge of probability and statistics. Further information can be found in Appendix A of the SSfM Best Practice Guide on Uncertainty and Statistical Modelling [12].

A **random variable** (rv) X is a quantity which can take any value in a prescribed range, the probability of the variable taking different values being determined by a distribution function $G_X(\xi)$. A random variable can be considered as a mapping from a space of events to the real line or some subset of it. For instance, if the space of events is “combinations of heads and tails that can occur in 16 tosses of a coin”, a random variable would be “number of heads”.

A **sample path** is the sequence of values a random variable takes during a simulation of its behaviour. The **transition probabilities** of a random variable are the probabilities that describe the distribution function of the values it could take next, given its current and previous values.

The **distribution function** $G_X(\xi)$ is defined as

$$G_X(\xi) = P(\omega \in \Omega : X(\omega) \leq \xi),$$

where ω is an event and Ω is the space of all possible events, which is the probability of an event occurring that leads to a value of X that is less than ξ . So, for the example of the random variable “number of heads in 16 tosses of a coin”, the distribution function would be

$$G_X(\xi) = \sum_{i=1}^{\xi} \frac{16!}{(16-i)!i!} p^i (1-p)^{\xi-i}$$

where p is the probability of getting a head on a single toss of the coin. The distribution function is sometimes called the cumulative distribution function (cdf). A joint distribution function of two random variables is

$$G_{XY}(\xi, \eta) = P(\omega \in \Omega : X(\omega) \leq \xi, Y(\omega) \leq \eta),$$

and this definition can be expanded to more than two variables. Two random variables X and Y are **independent** if their joint distribution function satisfies

$$G_{XY}(\xi, \eta) = G_X(\xi)G_Y(\eta)$$

where G_X and G_Y are the distribution functions of X and Y respectively. If two variables are not independent then they are **correlated**. Correlation is discussed further during the discussion of moments of random variables below.

Another useful concept is the **probability density function** (pdf), $g_X(\xi)$, which is defined as

$$g_X(\xi) = \frac{d}{d\xi} G_X(\xi).$$

This function only exists for continuous random variables, the nearest equivalent for discrete random variables being the probability function $P(X = \xi)$. As an example, a continuous random variable that is uniformly distributed on the interval $[0,1]$ would have distribution and density functions

$$G_X(\xi) = \begin{cases} 0, & \xi < 0, \\ \xi, & 0 \leq \xi \leq 1, \\ 1, & \xi > 1, \end{cases} \quad g_X(\xi) = \begin{cases} 0, & \xi < 0, \\ 1, & 0 \leq \xi \leq 1, \\ 0, & \xi > 1. \end{cases}$$

A **coverage interval** is an interval for which it can be stated with a given level of probability that the interval contains a specified proportion p of the population. If an interval at the $100p\%$ level ($0 < p < 1$) is required and the distribution function F_X is known, any interval of the form

$$[G_X^{-1}(q), G_X^{-1}(q + p)], \quad 0 \leq q \leq 1 - p$$

is a suitable interval, with a common choice being $q = (1 - p)/2$ as this gives a probabilistically symmetric interval.

The **expectation** or mean of a random variable X is given by

$$\mu = E(X) = \begin{cases} \sum_i x_i P(X = x_i) & \text{discrete case} \\ \int_{-\infty}^{\infty} \xi f_X(\xi) d\xi & \text{continuous case} \end{cases}$$

The n^{th} **moment** of X is given by $E(X^n)$ and the n^{th} **centred moment** is $E((X - \mu)^n)$. These higher moments measure the scatter and “shape” of the distribution about the mean. In particular the **variance** is given by

$$\sigma^2 = \text{Var}(X) = E((X - \mu)^2),$$

and the **standard deviation** is σ . The **covariance** of two random variables X and Y is defined as

$$\text{cov}(X, Y) = E((X - \mu_X)(Y - \mu_Y)) = E(XY) - \mu_X \mu_Y,$$

where μ_X and μ_Y are the means of X and Y respectively. If the two variables are independent, then

$$E(XY) = E(X)E(Y) = \mu_X \mu_Y \Rightarrow \text{cov}(X, Y) = 0.$$

If σ_X and σ_Y are the standard deviations of X and Y , then the **statistical correlation** of X

and Y is given by

$$\text{cor}(X, Y) = \text{cov}(X, Y) / \sigma_X \sigma_Y .$$

Two sets of samples from random variables X and Y can also have a **rank correlation**. Suppose that the values of the two variables are $X_1, X_2, \dots, X_N$, and $Y_1, Y_2, \dots, Y_N$. Each value is assigned a rank ordering using a function R such that

$$R(X_i) < R(X_j) \Rightarrow X_i < X_j$$

$$R(X_i) = R(X_j) \Rightarrow X_i = X_j$$

$$R(\min\{X_i, i = 1, 2, \dots, N\}) = 1$$

and R takes consecutive integer values unless two (or more) values are equal, in which case an averaged value of consecutive integers is used. The rank correlation coefficient, an approximation to the correlation coefficient defined above is given by

$$1 - 6 \sum_{i=1}^N \frac{(R(X_i) - R(Y_i))^2}{N(N-1)(N+1)} .$$

In metrology, **noise** can be described as any part of the signal that cannot be described by a deterministic model, and is usually caused by random uncontrolled factors in an experiment, such as temperature fluctuations, electro-magnetic field effects, or humidity variations. **White noise** is so called because it has a constant power spectrum over all frequencies, and can be modelled as a Wiener process as explained in section 5.1.2. Noise that is not white is called **coloured noise**.

3 Methods of uncertainty evaluation

Inclusion of random and uncertain effects in differential equations falls into two broad classes: one class in which the solution has a smooth path, and one where the path is not smooth and not differentiable.

If the random effect is included as the random choice of an input quantity or as a randomly forced initial condition then the solution has a smooth path and the problem is a random differentiable equation. Each realization (i.e. each pick of the parameter from a random set) can be solved using any of the standard deterministic methods. The methods used to choose suitable values of the input quantities are described in section 4 as “sampling methods”.

Sampling methods involve repeated runs of the model using particular choices for the values of the input quantities to construct the pdf of the output quantity by accumulation of results. They do not require the source code to be available, nor do they need details of the equations involved. This makes them suitable for use with black box software. However, they often require a large number of software runs to accumulate a detailed distribution, which can be a problem if a large model with a lengthy run time is being used.

If all the uncertainty in the system as a single noise term the solution will have a non-smooth, non-differentiable path. The problem becomes an ordinary differential equation with probabilistic forcing, a stochastic differential equation (SDE). There are a number of analytical methods that can be used to explore such problems, and these are described in section 5. In general, the analytical methods require information about the equations that make up the model and access to the source code and so are unlikely to be suitable for use with black box software. Often they involve assumptions about the pdfs of the input quantities and so may not be sufficiently general for some applications.

The decision of which type of method to use is partly determined by model complexity. If the model equations are highly nonlinear or very complicated, or if the model equations are unknown (for instance due to use of black box software), it is probably best to use a sampling method since these require fewest assumptions about the model and its input quantities. If the model is simple and linear and the input quantities have well-understood distributions, it is worth considering an analytic method as it may be more efficient in terms of time and may produce more information than the sampling method.

Another factor affecting choice of method is the final aim of the model. If a coverage interval is required, knowledge of the distribution function will be required in order to define the interval. However, it is possible that the uncertainty is to be used in a larger uncertainty budget, and that may mean that only a magnitude is necessary. For instance, the model output quantity may have an uncertainty associated with it that is significantly smaller than that from other sources, so a standard deviation and an assumption of normality may be sufficiently accurate for the purpose.

4 Sampling methods

Sampling methods treat the continuous model solution procedure as a black box. They take values randomly sampled from each of the pdfs of the input quantities, run the model with the sampled values as input quantities, and use the model results to get information about the joint pdf of the quantities of interest. Sampling methods do not involve propagation or convolution of distributions analytically, so they can accept more general pdfs than some analytic methods. Sampling methods are also widely applied to the calculation of complicated multi-dimensional integrals [27, Chapter 7], and much of their development has been driven by research in that area.

In general, sampling methods are straightforward to understand and implement, and are robust in as much as they can be applied to any problem involving a model that has input quantities and output quantities. Their main drawback, as will be discussed below, is that they generally require a large number of model evaluations to produce accurate and reliable information.

Many sampling methods exist that potentially could be applied to uncertainty evaluation in continuous models, but the three most commonly-used ones are described in the following sections. Examples of other methods include bootstrap sampling, quasi-Monte Carlo simulation, and Markov Chain Monte Carlo simulation. Details of some advanced methods are given as integration procedures in *Numerical Recipes* [27] in the chapter on random numbers.

It should be noted that because it is likely that the ideas underpinning these methods (in particular Monte Carlo simulation) will be familiar from evaluation of uncertainty in discrete modelling, no guidance is given on fundamentals such as how to sample from a distribution. Readers who are new to the subject are referred to the Best Practice Guide on Uncertainty and Statistical Modelling [12]. Most of the methods described here are also explained in *Sensitivity Analysis* [22], particularly in chapter 6.

4.1 Design of experiments and “extreme value” methods

Design of experiments and “extreme value” methods are used to try and obtain information about the output quantity’s limiting values in the minimum amount of time. In their simplest form, they rely on one key assumption: that the maximum and minimum values of the output quantity occur at maxima or minima of the input quantities. Sometimes their results are regarded as coverage intervals: this is not the case, since the methods generate no statistical information that could lead to a coverage interval. The main reason they are included in this report is to point out their failings as tools for uncertainty evaluation.

Extreme value methods are the simplest type of sensitivity test. They should not be used if at all possible as they assume linearity of the model and independence of the input quantities, and provide very little useful information. Suppose the model is $Y = f(\lambda_1, \lambda_2, \lambda_3, \dots, \lambda_n)$ where Y is the output quantity and the λ_i are the n input quantities. Suppose that the expected values of the parameters are $\{\mu_i, i = 1, \dots, n\}$, and that known limits exist for each parameter so that $\alpha_i \leq \lambda_i \leq \beta_i, i = 1, \dots, n$. Then an extreme values test would consist of $2n + 1$ model runs producing results y_i where

$$\begin{aligned}
y_1 &= f(\mu_1, \mu_2, \dots, \mu_n) \\
y_2 &= f(\alpha_1, \mu_2, \dots, \mu_n) \\
y_3 &= f(\beta_1, \mu_2, \dots, \mu_n) \\
&\vdots \\
y_{2j} &= f(\mu_1, \mu_2, \dots, \alpha_j, \dots, \mu_n) \\
y_{2j+1} &= f(\mu_1, \mu_2, \dots, \beta_j, \dots, \mu_n) \\
&\vdots \\
y_{2n} &= f(\mu_1, \mu_2, \dots, \alpha_n) \\
y_{2n+1} &= f(\mu_1, \mu_2, \dots, \beta_n)
\end{aligned}$$

It is then assumed that if $y_+ = \max\{y_i, i = 1, 2, \dots, 2n+1\}$ and $y_- = \min\{y_i, i = 1, \dots, 2n+1\}$ then $y_- \leq y \leq y_+$ for all possible values of the λ_i . This clearly neglects the effects of interaction between parameters, and assumes the relationship between the output quantity and the input quantities is monotonically increasing or decreasing. In many models these assumptions are invalid, as is illustrated in the example given below.

Design of experiments (DOE) methods were originally developed to minimise the number of experiments necessary to investigate the effects of varying the input quantities that could affect the output quantity. They are particularly common in quality processes, process control, and manufacturing and engineering applications where the input quantities are constrained to take values within some known range. They are similar in construction to the extreme values tests outlined above but their choice of parameter value combinations is slightly more sophisticated.

For a system with n parameters, an m -level DOE method takes m evenly spaced values from the range of each parameter, constructs all m^n possible sets of n values, and then chooses a subgroup of these sets for use in the model. The subgroup is usually chosen such that the effects of varying each of the individual parameters can be calculated from the results, and often such that the effects of interactions between pairs of values can be calculated as well. Details of the process of choosing parameter values and analysing results can be found in most quality control manuals [28], but generally the analysis of results is not used and the selection of tests is used to find maximum and minimum values for a coverage interval.

The main benefit of DOE methods is that they require far fewer runs to obtain information than other sampling methods. However, they must be used with extreme caution. The intervals produced do not correspond to a stipulated coverage probability and so are not “uncertainty estimates”. If all the input quantities used are assigned pdfs that are zero outside some closed interval and the output quantity depends linearly on them, then the interval formed from the maximum and minimum results of all possible combinations is a 100% coverage interval for the result, but this is a special case.

The other main drawbacks are that the method is very unlikely to work for non-linear models, and it is unlikely to give an accurate coverage interval if more than one input quantity is used. The coverage interval is likely to be an over-estimate, which may be acceptable for safe statement of a value but is generally not ideal for uncertainty evaluation since the aim in most metrology applications is to obtain a minimal interval that will include the output quantity value with a stipulated probability.

For example, consider the advection equation

$$\frac{\partial u}{\partial t} + \frac{\partial u}{\partial x} = 0, \quad 0 \leq x, t < \infty,$$

$$u(x, 0) = \lambda \exp\left[-(x - x^*)^2\right].$$

This has solution

$$u(x, t) = \lambda \exp\left[-(x - t - x^*)^2\right],$$

a wave travelling in the direction of increasing x . Suppose that the solution at $x = 10$, $t = 5$ is of interest. If the parameter λ is uncertain with a coverage interval $[\lambda_-, \lambda_+]$ at some level of probability and the parameter x^* is held fixed, then, because the model is linear in λ , using the extreme values of λ in the model will produce a coverage interval at the same level of probability for $u(10, 5)$ and so an extreme value test produces an acceptable result. However, if x^* is taken to be uncertain, the method will not be suitable. Consider the case where x^* is uniformly distributed between 3.4 and 6.4, and λ is uniformly distributed between 0.8 and 1.2. Taking the extreme values would give an interval $[0.06, 0.28]$ for $u(10, 5)$, but the maximum value is 1.2, occurring at $\lambda = 1.2$ and $x^* = 5$. The full response surface is shown in figure 2 and it is clear that the central ridge along the line $x^* = 5$ leads to extreme values being inadequate for a full description of the response.

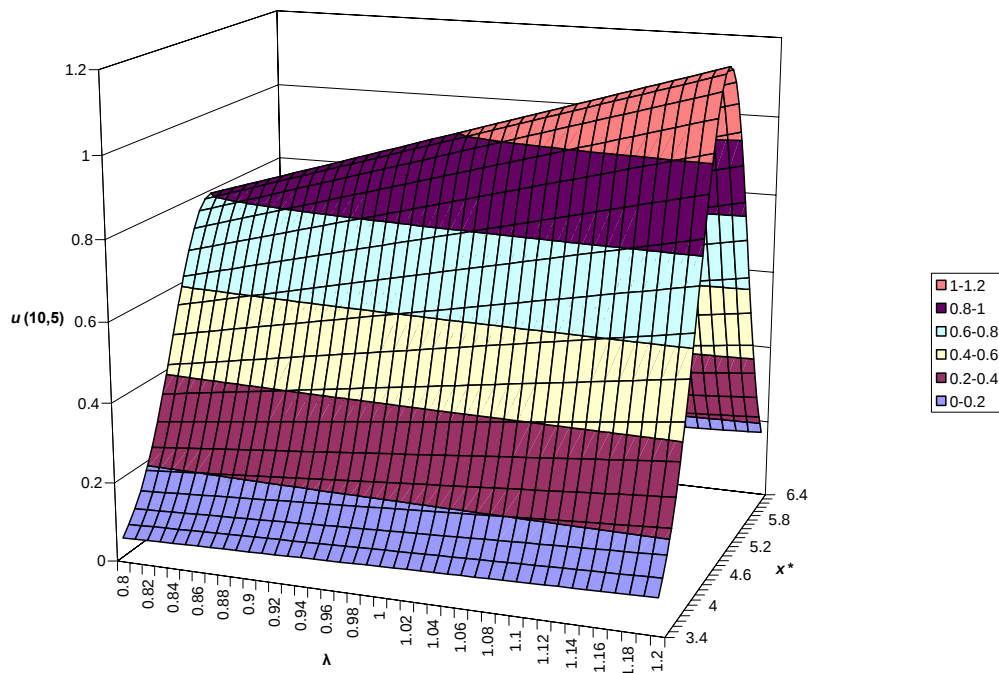


Figure 2: Plot of $u(10, 5)$ for various λ and x^* showing the central ridge that makes DOE methods unsuitable.

Whilst this example was clearly unlikely to work from inspection of the analytic solution, there are many other models where the dependence of the solution of interest on the initial and boundary conditions is not clear and so these methods should be avoided unless it is absolutely certain that the output quantity is linearly dependent on the input quantities.

4.2 Monte Carlo simulation

Monte Carlo simulation (MCS) is based on approximating the pdf or distribution function of the output quantity by running repeated model trials with input values sampled randomly from the joint pdf of the input quantities. MCS is a flexible approach to uncertainty evaluation that makes no extra assumptions about the problem beyond assignment of distributions for each of the uncertain input quantities. However, the method is such that these distributions can be of a very general form so the requirement is only that which is always needed for uncertainty quantification. MCS is becoming increasingly popular for evaluation of uncertainties in discrete models due to its flexibility and generality.

Suppose that the model for the quantity of concern is $Y = f(X_1, X_2, \dots, X_n)$ where $\mathbf{X} = \{X_i, i = 1, 2, \dots, n\}$ are the uncertain input quantities and f could be any model, from a simple analytic relationship to a complicated model using black box software. The MCS method can be viewed as a four step process:

1. Generate a sample vector of input quantity values $\mathbf{x}$ by sampling randomly from the joint pdf of the X_i . This joint pdf will be the product of the individual pdfs if the X_i are independent. Repeat this sampling M times to produce $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_M$.
2. Run the model M times to calculate $y_1 = f(\mathbf{x}_1), y_2 = f(\mathbf{x}_2), \dots, y_M = f(\mathbf{x}_M)$.
3. Construct an approximation to the distribution function of Y from the M values $y_1, y_2, \dots, y_M$.
4. Use the distribution function and the results to calculate any required statistical quantities.

The steps in this process that are most likely to be unfamiliar are the construction of the distribution function of Y from the results, and the choice of the number of trials M . These points are discussed further in sections 4.2.1 and 4.2.2.

The discussion above has assumed a single output quantity of interest Y . It is also possible to apply MCS to a situation where more than one quantity is of interest, as described in the SSfM Best Practice Guide (BPG) on Uncertainty and Statistical Modelling [12]. In that case, the process model is $\mathbf{Y} = \mathbf{f}(\mathbf{X})$ and $\mathbf{Y}$ is a vector with m components. The method as outlined using M samples above will produce a set of M vectors $\mathbf{y}_1, \mathbf{y}_2, \dots, \mathbf{y}_M$. The means of the components of $\mathbf{Y}$ can be estimated from the M samples, and subtracted from each of the samples to produce a set of M vectors, $\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_M$ say, whose m components all have mean zero. These vectors can be assembled into an m by M matrix, $\mathbf{Z} = \{\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_M\}$. Then an estimate of the covariance matrix is

$$V_y = \frac{1}{M-1} \mathbf{Z} \mathbf{Z}^T.$$

It is not simple to define an equivalent to the coverage interval for a multivariate output quantity. The problem is that an infinite number of shapes could be defined for a given coverage probability, making a sufficiently general definition impractical. However, a suggestion for a practical version of a coverage region is made in the BPG [12] that will be repeated here. For linear problems, the covariance matrix V_y defines a one standard-deviation ellipsoid, E_0 say, centred on the point denoting the estimate of $\mathbf{Y}$ (the ellipsoid definition is in the form of a scaling and rotation of a sphere, details are given in Mardia, Bibby and Kent [29]). If it is assumed that $\mathbf{Y}$ has a multi-dimensional Gaussian

distribution, the volume of an ellipsoid containing 95% (say) of the possible output quantities can be calculated. Then the ellipsoid with this volume that is concentric with E_0 will be a 95% coverage region for $\mathbf{Y}$. Two points should be emphasised: one is that this is only one of an infinite number of such regions, and the other is that this choice is dependent on the assumption that $\mathbf{Y}$ has a multivariate Gaussian distribution. If an alternative method for calculation of the volume of the ellipsoid is available that is more suitable for a given distribution, this should be used instead. Additionally, it should be noted that the ellipsoid might not be suitable for all multivariate distributions. However, it should be a reasonable initial guess for most smooth distributions.

The main drawback of applying MCS to continuous models is the large number of trials that are needed to obtain a good approximation to the pdf of the output quantity. This is particularly problematic for continuous models because often a continuous model that is sufficiently detailed to provide accurate results can have a run time of many hours. This makes MCS infeasible for large models, some transient models, and many non-linear models.

4.2.1 Constructing the distribution function

Suppose that M trials have been run, and the results are $y_1, y_2, \dots, y_M$. Then the steps in constructing an approximation to the distribution function are:

1. Reorder the results so that $y_1 \leq y_2 \leq \dots \leq y_M$.
2. Assign a cumulative probability $p_r = (r - 1/2)/M$ to y_r for $r = 1, 2, \dots, M$.
3. Construct the piecewise-linear function joining the points (y_r, p_r) , $r = 1, \dots, M$,
so that $\hat{G}(y) = p_r + \frac{y - y_r}{M(y_{r+1} - y_r)}$, $y_r \leq y \leq y_{r+1}$, $1 \leq r \leq M - 1$ is the
approximation to the density function $G_Y(y)$.

This will give a distribution function that is valid for $1/(2M) \leq p \leq 1 - 1/(2M)$, but it is likely that the distribution will be unreliable near to the endpoints of this interval. A typical distribution function is given in the example below.

The unsorted data can be used to construct an approximation to a pdf if required, by splitting the range of values that the y_i take into a number of bins of various widths, assigning each of the y_i to the appropriate bin, and creating a histogram. The smaller the bin width of the histogram, the more detailed the pdf will be. It is not recommended that the histogram be used for anything other than a visual inspection, for instance to check for symmetry, to check the number of peaks, or to see if the pdf resembles any of the common distributions. Some of the choices to be made during the construction of the histogram (e.g. number of bins, bin width) are subjective making it an unreliable tool. The distribution function will be more smooth than the pdf, as it is the integral of the pdf, which makes the distribution function a better choice for estimation of values such as coverage interval limits.

For example, the histogram shown in figure 3 is derived from the distribution function calculated by the MCS trials as part of the first case study (see section 7.1). It was thought that the distribution function looked like the output quantity had a triangular distribution, and so the pdf was constructed to check this theory. As the figure shows, a visual inspection of the pdf does not disprove this hypothesis.

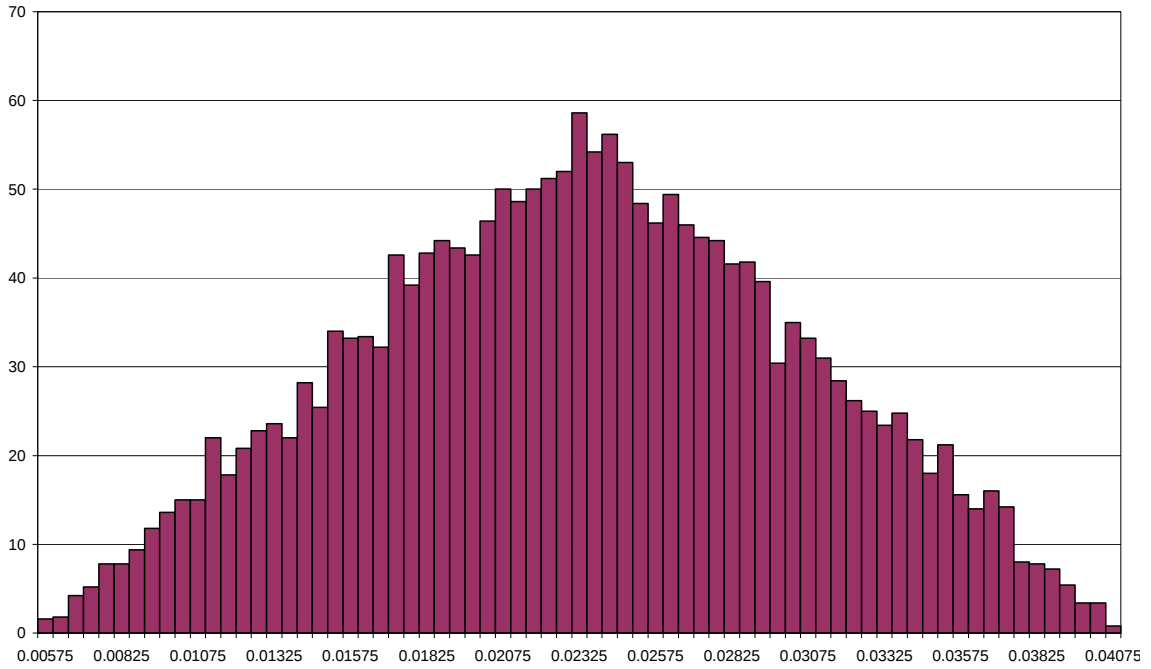


Figure 3: A histogram showing a pdf derived from the MCS results generated during the first case study (see section 7.1), showing an approximately triangular distribution.

The test results can be used to calculate a mean and standard deviation in the usual way:

$$\bar{y} = \frac{1}{M} \sum_{i=1}^M y_i, \quad s = \left(\frac{1}{M-1} \sum_{i=1}^M \{y_i - \bar{y}\}^2 \right)^{1/2}$$

These values are used as the value of the output quantity and the standard uncertainty associated with y respectively. The values are not identical to the expectation and the standard deviation of the distribution function, but for large values of M the difference will be negligible. The distribution function can be used to calculate coverage intervals. Suppose that a coverage interval corresponding to a $100p\%$ coverage probability is required, where p is a probability between 0 and 1, and that q is some value $0 \leq q \leq \min(p, 1-p)$. Then any interval of the form

$$[\hat{G}^{-1}(q), \hat{G}^{-1}(p+q)]$$

is a coverage interval at the $100p^{\text{th}}$ level, and the probabilistically symmetric coverage interval will have $q = (1-p)/2$.

The Supplemental Guide to the GUM [11] suggests that if the number of trials is restricted to M where M is less than 100 say, an alternative approach to that described above for defining the output quantity's pdf should be employed. If the trials have led to a mean value y and standard deviation s for the output quantity of interest, the guide suggests that the output quantity should be assumed to have distribution $y + Ts$ where T denotes the t -distribution with $M-1$ degrees of freedom. This is a smooth distribution with the same mean and standard deviation as the sample. This approach is compared with the method outlined above in the example in section 4.2.3.

4.2.2 Choosing M

The choice of M is a compromise between carrying out enough model runs that the results will provide a good approximation to the distribution of Y and carrying out as

few runs as possible in order to obtain the approximation economically. How good an approximation to the distribution is needed depends on which statistical quantities are to be calculated. More trials may be required for calculation of coverage intervals than for calculation of mean and variance.

Some advice for choosing M exists. The supplement to the GUM [11] suggests that 10^6 model runs can often be expected to provide a 95% coverage interval that is accurate to one or two significant decimal digits. The authors of this report have found that 10 000 trials can be sufficient in some cases (as will be shown in the case studies).

The supplement to the GUM [11] also describes an adaptive procedure for choosing M . It is based on carrying out an increasing number of Monte Carlo trials until the quantities of interest have stabilised in a statistical sense. The practical approach described in the supplement can be summarised as follows:

1. Calculate a sequence of M Monte Carlo trials to obtain output quantities $y_i, i = 1, 2, \dots, M$.
2. For each sequence of output quantities $y_i, i = 1, 2, \dots, h \leq M$, calculate the mean, the standard uncertainty, and the endpoints of a 95% coverage interval. Denote these quantities by $y^{(h)}, u(y^{(h)}), y_{\text{low}}^{(h)}$, and $y_{\text{high}}^{(h)}$.
3. Calculate the arithmetic mean and standard deviation of each of the sequences of quantities $\{y^{(h)}, h = 1, 2, \dots, M\}$, $\{u(y^{(h)}), h = 1, 2, \dots, M\}$, $\{y_{\text{low}}^{(h)}, h = 1, 2, \dots, M\}$, and $\{y_{\text{high}}^{(h)}, h = 1, 2, \dots, M\}$. Denote the standard deviations by $s_y, s_{u(y)}, s_{y_{\text{low}}}$ and $s_{y_{\text{high}}}$.
4. Compare the required degree of approximation in $u(y)$ to the largest of $2s_y, 2s_{u(y)}, 2s_{y_{\text{low}}}$ and $2s_{y_{\text{high}}}$. If the required degree of approximation is greater than double the largest standard deviation, the overall computation is regarded as having stabilised. The results from the total number of MC trials are then used to provide the output quantity estimate, the associated standard uncertainty, and the coverage interval for the output quantity value.

This procedure will be used in the example below in order to check the statistical stability of the results.

4.2.3 An example of MCS

Consider the advection equation

$$\begin{aligned} \frac{\partial u}{\partial t} + \frac{\partial u}{\partial x} &= 0, \quad 0 \leq x, t < \infty, \\ u(x, 0) &= \lambda \exp\left[-(x - x^*)^2\right], \end{aligned} \tag{1}$$

as before, but now assume that the result at $x = 1.0, t = 0.5$ is of interest, and that x^* is uniformly distributed between 0.2 and 0.8, and λ is uniformly distributed between 0.9 and 1.1 independently of x^* . A finite difference model is used to solve (1) (validated against the exact solution), and a Monte Carlo simulation is run using 25 000 trials.

As was stated in section 4.1, the problem (1) has an analytic solution $u(x, t) = \lambda \exp[-(x - t - x^*)^2]$ so the analytic expression for the output quantity is $y = \lambda \exp[-(0.5 - x^*)^2]$. Since λ and x^* are independently distributed, their joint pdf is

$$g_{x^* \Lambda}(x^*, \lambda) = \frac{1}{(0.8 - 0.2)(1.1 - 0.9)} = \frac{1}{0.12},$$

and so the expressions for the mean and variance of y are

$$\begin{aligned} \mu &= \int_{0.9}^{1.1} \int_{0.2}^{0.8} \frac{\lambda \exp[-(0.5 - x^*)^2]}{0.12} dx^* d\lambda = \left[\frac{\lambda^2}{2} \right]_{0.9}^{1.1} \int_{-0.3}^{0.3} \frac{\exp[-t^2]}{0.12} dt \\ &= 0.2 \frac{2}{0.12} \int_0^{0.3} \exp[-t^2] dt \approx 0.9708 \end{aligned}$$

$$\begin{aligned} \sigma^2 &= \int_{0.9}^{1.1} \int_{0.2}^{0.8} \frac{1}{0.12} (\lambda \exp[-(x - x^*)^2] - \mu)^2 dx^* d\lambda \\ &= \frac{1}{0.12} \int_{0.9}^{1.1} \int_{0.2}^{0.8} \lambda^2 \exp[-2(x - x^*)^2] dx^* d\lambda - \mu^2 \\ &= \frac{1}{0.12} \left[\frac{\lambda^3}{3} \right]_{0.9}^{1.1} \int_{0.2}^{0.8} \exp[-2(x - x^*)^2] dx^* - \mu^2 \\ &= \frac{1}{0.12} \left[\frac{1.1^3 - 0.9^3}{3} \right] 2 \int_0^{0.3} \exp[-2t^2] dt - \mu^2 \approx 0.0617 \end{aligned}$$

Figure 4 shows the calculated distribution functions at 25, 250, 2500 and 25 000 trials, and table 1 shows the mean, standard deviation, and 2.5th and 97.5th percentiles (for a 95% coverage interval), calculated from the results.

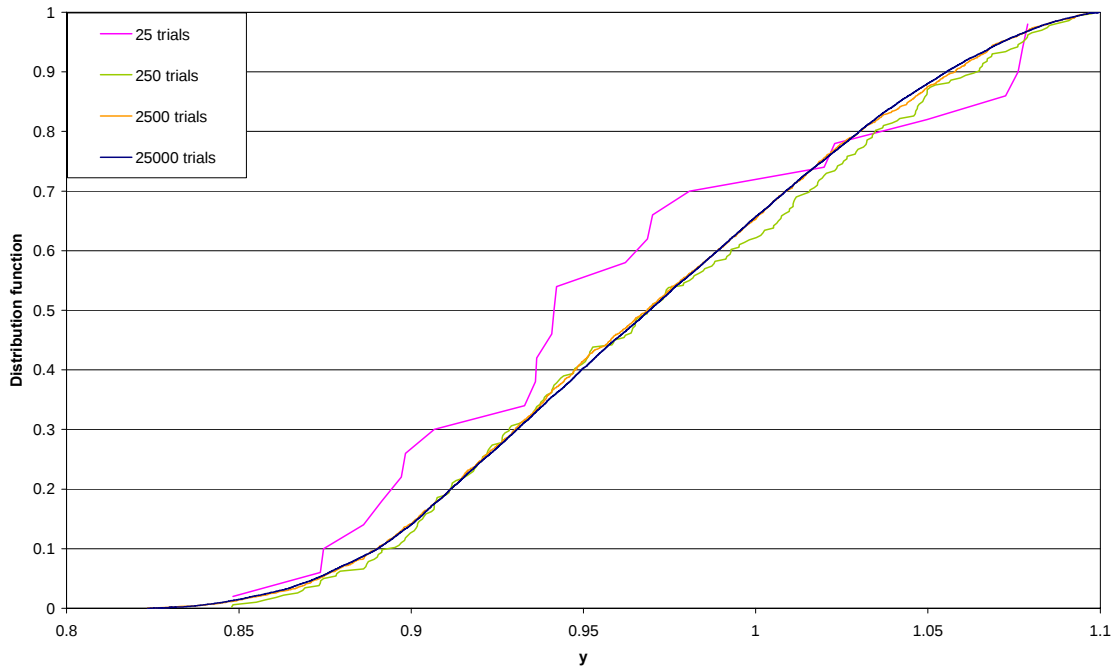


Figure 4: Calculated distribution functions based on 25, 250, 2500 and 25 000 trial runs.

Number of trials	Sample mean	Sample standard deviation	2.5 th percentile	97.5 th percentile
25	0.959	0.071	0.852	1.079
250	0.973	0.063	0.866	1.085
2 500	0.970	0.062	0.860	1.081
25 000	0.970	0.062	0.858	1.082

Table 1: Statistical quantities calculates from the MCS results with different numbers of trials.

It is clear that as the number of trials increases, the approximations to the distribution function, the mean, and the standard deviation become more stable. Calculating the same quantities using the analytic solution for the values λ of and x^* generated during the 25 000 trial run gave a mean of 0.970 and a standard deviation of 0.062, so the differences between the mean and standard deviation calculated from the MCS results and the “true” values calculated above are largely due to the use of MCS rather than the use of a finite difference approximation in the implementation of the model.

Figure 5 shows the distribution function based on 25 000 trials compared with the approximation suggested in section 4.2.1 for use with a small number of trials, based on the first 25 trials. Whilst the approximation is not a particularly good one to the 25 000 trial results, it is an improvement over the distribution function produced by linear interpolation between the 25 trial points by virtue of its regularity and smoothness. A 95% coverage interval from the t -distribution approximation would be [0.813, 1.106]. This is a larger interval than those calculated from MCS results, which is to be expected from an estimate based on fewer data points.

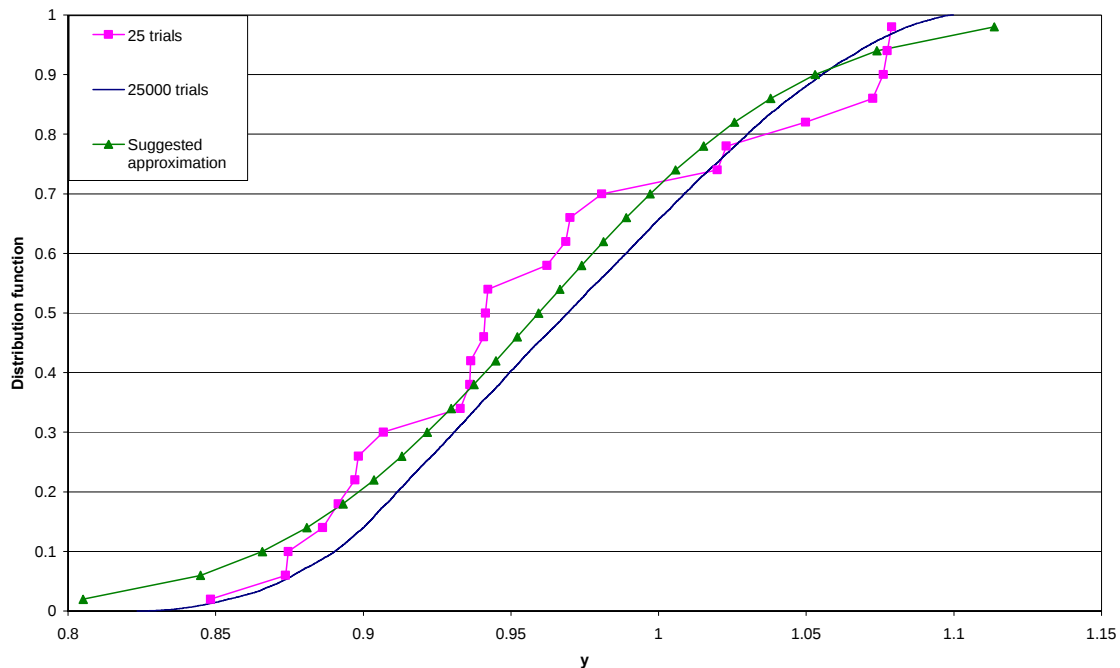


Figure 5: A plot of the distribution functions based on 25 and 25 000 trials, and the approximation to the 25 trial distribution as suggested in section 4.2.1.

Applying the results on coverage intervals given in section 4.2.2 to the tests gives the results shown in table 2 for the four standard deviations s_y , $s_{u(y)}$, $s_{y\text{low}}$ and $s_{y\text{high}}$ calculated after different numbers of trials. If an uncertainty accurate to two figures was required, it is likely that a few more tests would be needed since this would require all of the standard deviations given below to be less than 0.005.

Number of trials	s_y	$s_{u(y)}$	$s_{y\text{low}}$	$s_{y\text{high}}$
25	0.1939	0.0205	0.1747	0.2158
250	0.0614	0.0073	0.0548	0.0685
2 500	0.0195	0.0025	0.0175	0.0217
25 000	0.0062	0.0009	0.0056	0.0069

Table 2: Standard deviations used in the test for statistical stability of the calculated coverage interval, for different numbers of MCS trials.

The increase in stability with increasing number of trials is clear from the figures in table 2. Figure 6 shows a log-log plot of the standard deviations against the number of trials. This shows that the convergence behaves like M^d where d is some constant. Straight-line fits to the data show that $d = 0.5$ for all four lines in this case.

It is also useful to note that the largest standard deviation is $s_{y\text{high}}$. In general, the end points of the coverage interval are expected to stabilise more slowly than the mean and the standard uncertainty since stable coverage interval limits require knowledge of the tails of the distribution function. This has not happened for this example: s_y is of the same order of magnitude as $s_{y\text{low}}$ and $s_{y\text{high}}$. It is not clear whether s_y is unusually large or whether $s_{y\text{low}}$ and $s_{y\text{high}}$ are unusually small.

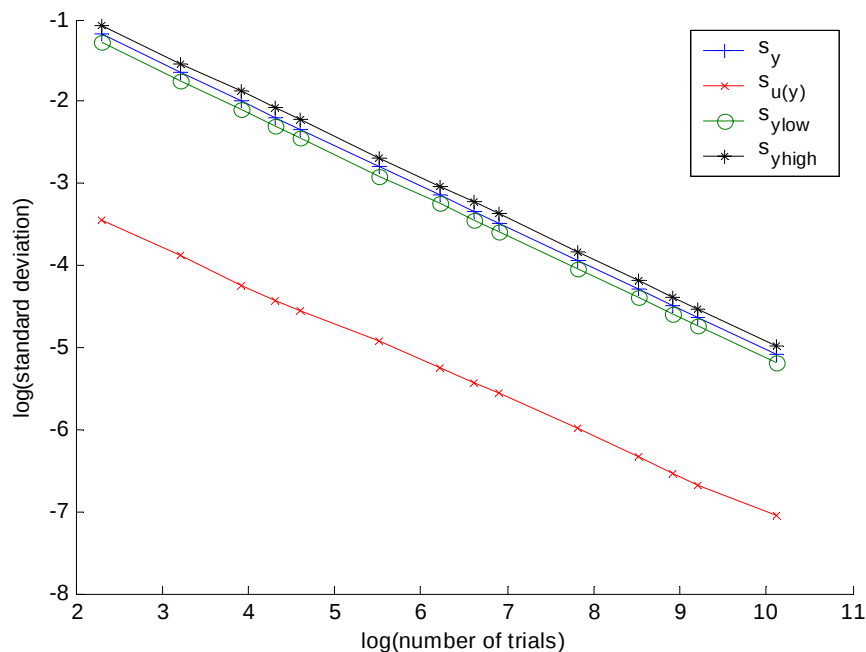


Figure 6: Plot of $\log(\text{number of trials})$ versus $\log(\text{standard deviation})$ for s_y , $s_{u(y)}$, $s_{y\text{low}}$ and $s_{y\text{high}}$, showing the convergence behaviour. The slopes of the best straight line fits are -0.50 for s_y , $s_{y\text{low}}$ and $s_{y\text{high}}$, and -0.46 for $s_{u(y)}$.

4.3 Constrained Monte Carlo simulation

As mentioned in section 4.2, the main drawback of MCS is the large number of trials necessary to obtain a reliable approximation to the output quantities' distribution function. This is a particular problem for continuous models as models commonly take several CPU hours to run, making MCS with 10 000 an impossible proposition. With this problem in mind, methods have been developed to obtain information about the pdf of the output quantity based on fewer model runs. Constrained Monte Carlo simulation methods use the same basic method as MCS, but constrain the sampling to build a reasonably accurate approximation to the distribution function more quickly. Since these methods use fewer model runs than the full MCS method, they cannot produce coverage intervals to the same accuracy as a full MCS, so they should be used with caution, but they can provide useful information quickly. Constrained MCS methods become equivalent to Monte Carlo simulation as the number of samples tends towards infinity.

Two types of method will be described in this section: Stratified sampling, and Latin Hypercube sampling. Stratified sampling is a fairly general term that covers a number of methods based on similar ideas, and Latin Hypercube sampling is probably the most extensively used stratified sampling technique. Finally the techniques will be compared with regular MCS using a worked example.

4.3.1 Stratified sampling

Stratified sampling [13] aims to reduce the number of model runs required whilst ensuring all areas of the distributions of the input quantities are covered. It achieves this by imposing restrictions on the samples to be used.

Suppose the problem of interest is expressed as

$$Y = f(\mathbf{X}); \quad \mathbf{X} \in \Omega \subset R^n,$$

where $\mathbf{X} = \{X_1, X_2, \dots, X_n\}$ are the inexact input quantities and Y is the output quantity of interest (taken to be a scalar quantity in this case, although this is not generally essential for application of the method). The procedure for stratified sampling is:

- Divide the domain, Ω , of the input quantities of the problem into N regions $\{S_i, i=1, 2, \dots, N\}$ such that no two regions overlap and the union of all the regions is the whole of the space Ω .
- Randomly sample m_i sets of values $\{\mathbf{x}^{ij}, j = 1, 2, \dots, m_i\}$ from each region S_i according to the joint pdf of the input quantities.
- Run the model with each set of parameters to produce a set of results y_{ij} .

It is clear that if $N = 1$, the method is the same as unconstrained MCS. Stratified sampling is sometimes called importance sampling, since the flexibility of the sampling method means that the inclusion in the sampling of regions of the input quantity space that are unlikely to occur in practice, but have an important effect when they do, can be ensured.

Suppose the probability $P(\mathbf{X} \in S_i) = p_i$, and $M = m_1 + m_2 + \dots + m_N$. It can be shown [13] that an unbiased estimator for the distribution function $P(Y \leq y)$ is given by

$$\bar{P}(y) = \sum_{i=1}^N \frac{p_i}{m_i} \sum_{j=1}^{m_i} g(y, y_{ij}),$$

$$g(y, y_{ij}) = \begin{cases} 1, & y - y_{ij} \geq 0, \\ 0, & y - y_{ij} < 0. \end{cases}$$

The unbiased estimators for the mean and variance are, respectively,

$$\bar{y} = \sum_{i=1}^N \frac{p_i}{m_i} \sum_{j=1}^{m_i} y_{ij},$$

$$s^2 = \sum_{i=1}^N \frac{p_i}{m_i} \sum_{j=1}^{m_i} (y_{ij} - \bar{y})^2.$$

If the m_i and the p_i are chosen so that $m_i = p_i M$, called proportional sampling, it can be shown [13] that the variances of the estimators are at least as small as those of the estimators generated by an MCS test with the same number of trials.

The most straightforward general implementation of this method divides the domain into regions of equal probability so that $p_i = 1/N$, $i = 1, 2, \dots, N$, but this may not be the best strategy if one particular region of the input space is of particular interest. If the input quantities are independent, a simple option is to divide all of the independent univariate distribution functions into k regions of equal probability by subdividing the vertical cumulative probability scale into equally-sized intervals and then using the distribution function to find the corresponding values of x_i , as illustrated in figure 7 for $k = 10$. This divides the space Ω into k^n cells of equal probability, so one sample can be taken from each cell. This strategy becomes less feasible as n increases, but it may be a good initial choice for small n where little is known about the dependence of y on the input quantities.

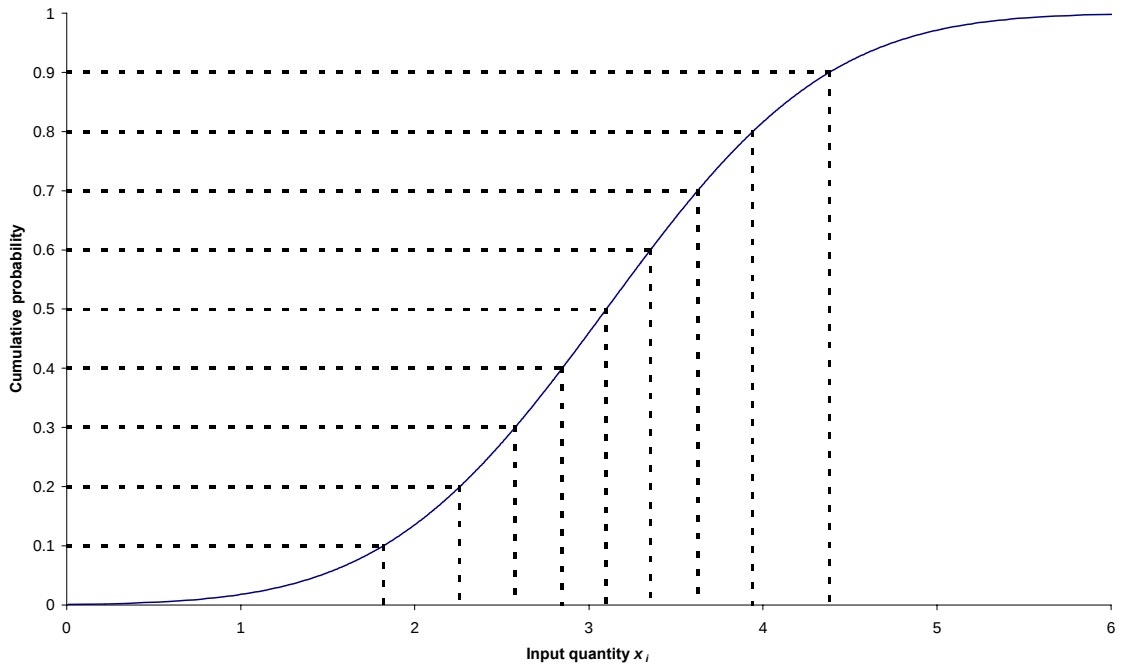


Figure 7: Subdivision of the space Ω into regions of equal probability by inversion of the distribution function.

The method has been extended to produce a class of methods called adaptive importance sampling methods [9, 27]. These methods use the results of one set of trials to generate the choice of samples used for the next set of trials. The version outlined in *Numerical Recipes* [27] combines stratified and adaptive sampling, allows the user to fix the maximum number of trials, and then calculates the optimal distribution of sample points adaptively. The adaptive calculation assumes that the pdf of the output quantity can be split into a product of functions of the input variables, and uses the calculation results to improve the approximation properties of these functions.

4.3.2 Latin Hypercube sampling

Latin Hypercube sampling (LHS) [13-15] is a form of stratified sampling method. Unlike the stratified sampling methods outlined above, the divisions of the domain from which the samples are taken do not have to cover the whole domain when put together.

For a model with n independent input quantities $X_1, X_2, \dots, X_n$, the method can be described in steps as follows:

1. Divide each of the n input spaces into k regions of equal probability by inverting the distribution function as shown in figure 7, and then take a random sample from each region for each pdf, giving a total of nk samples $\{x_{ij}, i = 1, 2, \dots, n, j = 1, 2, \dots, k\}$. It is required that $k \geq n$.
2. Construct n random permutations of the integers $1, 2, \dots, k$. Suppose that the j^{th} number in the i^{th} permutation is $q_i(j)$ $i = 1, 2, \dots, n, j = 1, 2, \dots, k$.
3. Assemble k sets of input quantity values so that the j^{th} set, $\mathbf{x}^j$ is $\mathbf{x}^j = \{x_{i,k}: q_i(j) = k, i = 1, 2, \dots, n\}$. These sets are the sets of input quantities to be used in the model.
4. Run the model k times to generate the output quantities $y_i, i = 1, 2, \dots, k$.
5. Calculate statistics and distribution function using the method as outlined in section 4.2.1.

As an illustration, consider a model with two input quantities, X_1 , which is uniformly distributed on the interval $[0, 1]$, and X_2 , which is independently uniformly distributed on the interval $[3, 5]$. Choose $k = 4$. Then a possible set of input samples is $x_{1,j} = \{0.136, 0.384, 0.592, 0.817\}$, $x_{2,j} = \{3.48, 3.92, 4.15, 4.71\}$. Suppose the two permutations of $(1, 2, 3, 4)$ are $(3, 2, 1, 4)$ and $(2, 1, 4, 3)$. Then the four pairs of input quantities are $\{0.592, 3.92\}$, $\{0.384, 3.48\}$, $\{0.136, 4.71\}$, $\{0.817, 4.15\}$. This is illustrated in figure 8.

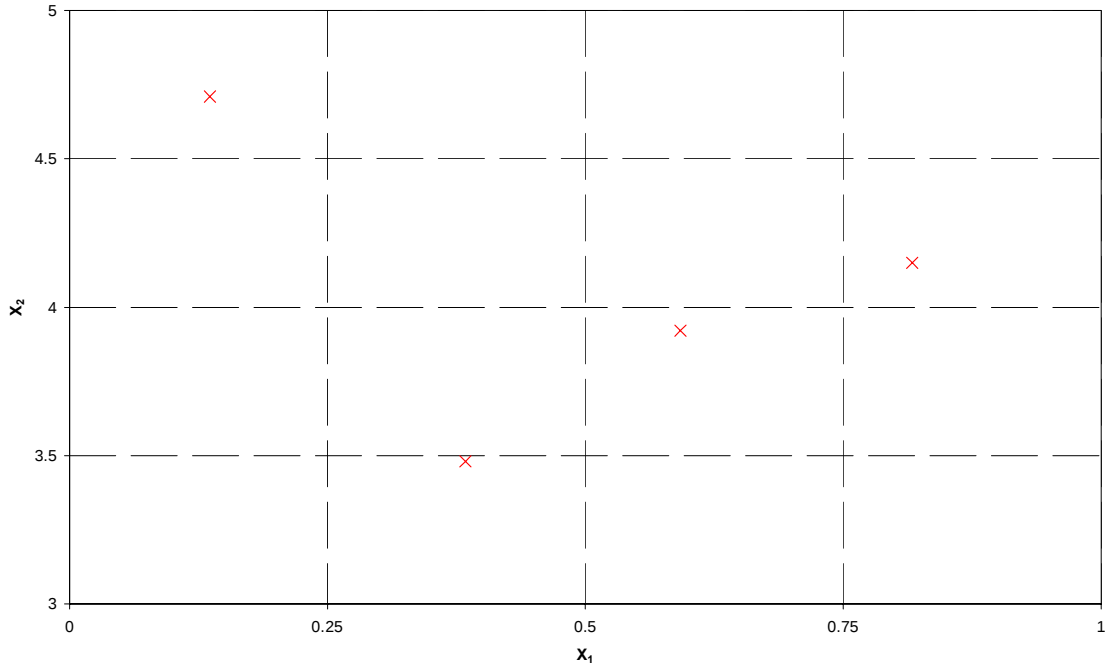


Figure 8: Illustration of the input quantity sets chosen with the LHS method. Chosen sets are denoted by red crosses. Note that there is one sample chosen from each row and column.

The method outlined above is only applicable to models with independent input quantities. There are two issues relating to covariance: the first is that the restriction on the sampling may have introduced spurious correlation that needs to be removed, and the second is that if the inputs are mutually dependent, that needs to be reflected in the sampling. It should be noted that throughout the following, the correlation that is used is the rank correlation rather than the sample correlation. The two are closely linked, but rank correlation is easier to manipulate in this case.

Iman and Conover [30] show that these problems can be dealt with by making changes to the permutations used to produce the samples. The technique has been compared with MCS and LHS without correlation control [14], and the following description is taken from that reference. Suppose that $\mathbf{P}$ is the k by n matrix with columns corresponding to the chosen permutations of $1, 2, \dots, k$, so for instance the matrix for the example in figure 8 would be

$$\mathbf{P} = \begin{bmatrix} 3 & 2 \\ 2 & 1 \\ 1 & 4 \\ 4 & 3 \end{bmatrix}.$$

Then the steps in the method are:

1. Divide the elements of $\mathbf{P}$ by $k + 1$, and map the resulting quantities (between 0 and 1) onto the Gaussian distribution with mean 0 and variance 1, so that $Z_{ij} = \Phi^{-1}(P_{ij}/(k+1))$, where Φ is the distribution function for the standardised Gaussian distribution.
2. Evaluate the covariance matrix of $\mathbf{Z}$ and form its Cholesky decomposition [31]

it into the form $\text{cov}(\mathbf{Z}) = \mathbf{L}\mathbf{L}^T$, where $\mathbf{L}$ is a lower triangular matrix

3. Calculate $\mathbf{Z}^* = \mathbf{Z}(\mathbf{L}^{-1})^T$.
4. Within each column of $\mathbf{Z}^*$, label the elements starting with 1 for the most negative element and increasing the label by 1 each time. So, for instance, if a column of $\mathbf{Z}^*$ was $\mathbf{z} = [0.3, -0.56, 0.28, -1.04, 1.2]^T$ the corresponding column of labels would be $[4 \ 2 \ 3 \ 1 \ 5]^T$ because $z_4 < z_2 < z_3 < z_1 < z_5$.

As an example, it can be shown that the process described above turns $\mathbf{P}$, which has strong correlation between the columns, into $\mathbf{P}^*$ which does not.

$$\mathbf{P} = \begin{bmatrix} 1 & 2 & 3 & 4 & 5 & 6 & 7 & 8 & 9 & 10 \\ 2 & 1 & 4 & 3 & 6 & 5 & 8 & 7 & 10 & 9 \end{bmatrix}^T, \quad \text{cov}(\mathbf{P}) = \begin{bmatrix} 9.1667 & 8.6111 \\ 8.6111 & 9.1667 \end{bmatrix}$$

$$\mathbf{P}^* = \begin{bmatrix} 2 & 1 & 4 & 3 & 6 & 5 & 8 & 7 & 10 & 9 \\ 1 & 6 & 3 & 7 & 4 & 8 & 5 & 9 & 2 & 10 \end{bmatrix}^T, \quad \text{cov}(\mathbf{P}^*) = \begin{bmatrix} 9.1667 & 1.722 \\ 1.722 & 9.1667 \end{bmatrix}$$

If a particular covariance between the inputs is required, this can be achieved by altering step 3 above to be “Calculate $\mathbf{Z}^* = \mathbf{Z}(\mathbf{L}^{-1})^T \mathbf{\Lambda}^T$ ”, where the desired covariance matrix has Cholesky factorisation $\mathbf{\Lambda}\mathbf{\Lambda}^T$. So for instance if the desired covariance matrix for the problem in figure 8 is

$$\text{cov}(\hat{\mathbf{Y}}) = \begin{bmatrix} 0.51 & 0.49 \\ 0.49 & 0.51 \end{bmatrix}, \quad \mathbf{\Lambda} = \begin{bmatrix} 0.7141 & 0.6861 \\ 0 & 0.1980 \end{bmatrix}$$

and so

$$\mathbf{Y} = \begin{bmatrix} -0.0544 & -0.0717 \\ -1.2703 & -0.2383 \\ 0.1806 & 0.2383 \\ 1.1440 & 0.0717 \end{bmatrix}, \quad \mathbf{P}^* = \begin{bmatrix} 2 & 2 \\ 1 & 1 \\ 3 & 4 \\ 4 & 3 \end{bmatrix}$$

is a better choice of permutation matrix. Illustrating with a more obvious example,

$$\mathbf{P} = \begin{bmatrix} 1 & 2 & 3 & 4 & 5 & 6 & 7 & 8 & 9 & 10 \\ 2 & 1 & 4 & 3 & 6 & 5 & 8 & 7 & 10 & 9 \end{bmatrix}^T, \quad \text{target cov}(\mathbf{P}) = \begin{bmatrix} 9 & -8 \\ -8 & 9 \end{bmatrix},$$

gave an improved permutation matrix of

$$\mathbf{P}^* = \begin{bmatrix} 9 & 10 & 7 & 8 & 5 & 6 & 3 & 4 & 1 & 2 \\ 2 & 1 & 4 & 3 & 6 & 5 & 8 & 7 & 10 & 9 \end{bmatrix}^T, \quad \text{cov}(\mathbf{P}^*) = \begin{bmatrix} 9.1667 & -9.1667 \\ -9.1667 & 9.1667 \end{bmatrix}$$

which, again, is a clear improvement.

McKay [13] shows that for a model with independent input quantities, the variances of the unbiased estimators for the mean and standard deviation are at least as small as those of the estimators generated by an MCS test with the same number of trials, provided that the model $Y = f(\mathbf{X})$ is monotonic in each of its inputs X_i .

An extension of this method, replicated LHS [15], can be used to investigate effects of individual input quantities on the result Y . The method can be used to form a quantity that could be viewed as a sensitivity coefficient: the conditional variance $\text{Var}[E[Y|X_i]]$. Replicated Latin Hypercube sampling takes k samples of the n input quantities as in

step 1 above, but then produces r replicate LHS tests using these samples. Each replicate is created by using a different random permutation matrix $\mathbf{P}$, so whilst all the replicates use the same values of the input quantities, they are grouped differently within each replicate, forming different $\mathbf{x}^j$ in step 3 above. McKay [15] gives a full explanation of the analysis of the results, but in brief the steps (based on ANOVA) are

- For each x_i , $i = 1, \dots, n$, assemble the sets $\{y_{jl}, l = 1, \dots, r\}$, $j = 1, \dots, k$ where y_{jl} is the result of the trial in the l^{th} replicate where x_i took its j^{th} value. This is just a reordering of the results.

- Calculate the following quantities:

$$\bar{y}_j = \frac{1}{r} \sum_{l=1}^r y_{jl}, \quad \bar{y} = \frac{1}{k} \sum_{j=1}^k \bar{y}_j,$$

$$SST = \sum_{j=1}^k \sum_{l=1}^r (y_{jl} - \bar{y})^2, \quad SSB = r \sum_{j=1}^k (\bar{y}_j - \bar{y})^2, \quad SSW = \sum_{j=1}^k \sum_{l=1}^r (y_{jl} - \bar{y}_j)^2$$

As a check, $SST = SSB + SSW$.

- The last three quantities are estimates of the total variance (SST), the between-sample variance (SSB) and the within-sample variance (SSW) for X_i .
- If Y is not sensitive to the value of X_i then all of the $\bar{y}_j$ should be approximately the same, so SSB will be small compared to SSW .
- Conversely, if Y depends only on X_i then for a fixed value of X_i , Y will be constant and so SSW will be small (in fact it will be zero) compared to SSB .
- Hence the ratio SSB/SST (which must lie between 0 and 1) can be used as an indication of how sensitive Y is to X_i .

This indicator may not be suitable for all models as it may be unduly influenced by outliers, but it is generally a good place to start if sensitivity is of interest.

4.3.3 An example of constrained MCS methods

The problem (1) was examined using Stratified sampling (SS) and Latin Hypercube sampling (LHS). Three different sets of tests were run: one with 25 trials, one with 50 trials, and one with 100 trials. Each test consisted of one SS test, one LHS test, and one MCS test. In all cases, the SS trial was carried out with regions of equal probability and a single sample per cell, so the implementation was proportional sampling.

In each test the mean and standard deviation were calculated, and these were compared with the analytic values. The test results were used to construct distribution functions, and the values of y and different levels of probability were compared with those taken from the 25 000 trial MCS test run previously (see section 4.2.3), which was regarded as being a good approximation to the correct solution. The distribution functions were compared visually and by calculating the root mean square difference between each test and the MCS test with 25000 trials by calculating

$$RMS = \sqrt{\frac{1}{M} \sum_{i=1}^M (y_i - z_i)^2},$$

where M is the number of test points, y_i is the result of the test and z_i is the result of the MCS test with 25 000 trials at the same level of cumulative probability.

The results of the calculations of means and standard deviations are shown in table 3. As in section 4.2.3, the analytical solutions are $\mu = 0.9708$, $\sigma = 0.0617$. These results show that for a small number of trials, LHS provides a better estimate of the mean than MCS and that SS provides a better estimate of the standard deviation than MCS.

To investigate these effects, 10 tests were run using SS and LHS with $M = 50$ and $M = 100$. The results of the repeated tests are shown in table 4. The table shows the mean and variance of the test means and standard deviations calculated for each method, i.e.

$$\text{Mean} = \frac{1}{10} \sum_{j=1}^{10} z_j \quad \text{Var} = \frac{1}{9} \sum_{j=1}^{10} (z_j - \text{Mean})^2$$

where z_j is the mean or standard deviation of test j . This shows the test-to-test variability of the methods.

M	Stratified sampling		LHS		MCS	
	μ	σ	μ	σ	μ	σ
25	0.966	0.063	0.970	0.055	0.959	0.071
50	0.970	0.062	0.971	0.061	0.966	0.064
100	0.972	0.062	0.971	0.058	0.967	0.070

Table 3: Means and standard deviations calculated from the test results.

M	Stratified sampling				LHS			
	Distribution mean		Distribution s.d.		Distribution mean		Distribution s.d.	
	Mean	Var	Mean	Var	Mean	Var	Mean	Var
50	0.9700	1.5×10^{-3}	0.0623	2.3×10^{-3}	0.9707	2.9×10^{-4}	0.0612	3.9×10^{-3}
100	0.9707	8.0×10^{-4}	0.0614	1.0×10^{-3}	0.9707	1.9×10^{-4}	0.0620	2.7×10^{-3}

Table 4: Means and standard deviations calculated from the results of repeated tests.

The results in table 4 support the hypothesis that LHS gives the better estimate of the mean (it is more accurate and has a lower variance) and that SS gives the better estimate of the standard deviation. Ideally this should be tested further since ten tests may not be sufficient, and it is probable that correlated input quantities will affect this behaviour, so general conclusions cannot necessarily be drawn. By comparison, a 500 trial MCS run gives a mean of 0.9693 and a standard deviation of 0.0625, and a 1000 trial MCS run gives a mean of 0.9697 and a standard deviation of 0.0614.

A typical set of distribution functions, calculated with the results of the 50 trial tests, are shown in figure 9. The figure shows that the MCS test with 50 trials deviates the most from the 25 000 trial MCS test. This is confirmed by the figures in table 5, which are the RMS differences calculated for each value of M and each method.

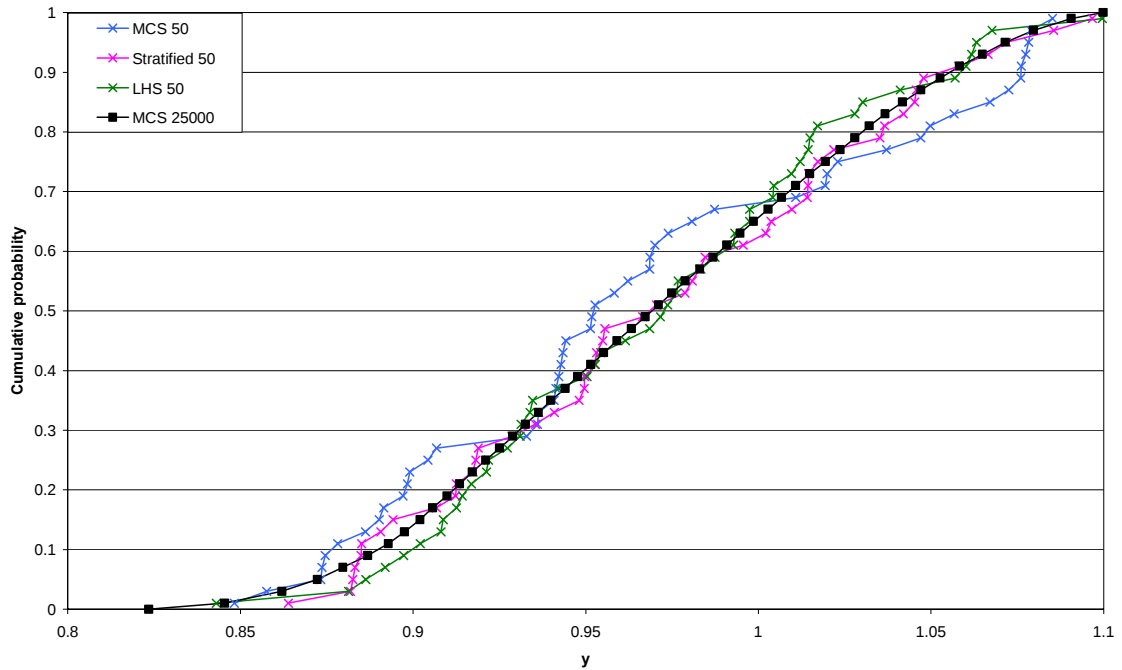


Figure 9: Typical distribution functions from Monte Carlo simulation (blue line), Stratified sampling (pink line), and Latin Hypercube sampling (green line). The black line shows values from the distribution function calculated using 25 000 MCS trials, which can be regarded as a good approximation to the correct solution.

M	Stratified sampling	LHS	MCS
25	8.0×10^{-3}	1.46×10^{-2}	1.95×10^{-2}
50	5.93×10^{-3}	7.09×10^{-3}	1.40×10^{-2}
100	3.16×10^{-3}	4.52×10^{-3}	6.88×10^{-3}

Table 5: RMS differences between the distribution functions generated by M -trial tests using three different methods and the results of the 25 000 trial MCS test sampled at appropriate levels of probability.

5 Analytical methods

As was mentioned in section 4, sampling methods treat the computational model as a black box and only consider the model's input and output quantities. In many cases, particularly for non-linear models, this is an advantage. However, for some models this neglects potentially useful information. The simplest example of this is to consider a linear model $Y = X_1 + X_2$ where X_1 and X_2 are independent Gaussian random variables. Then Y will also be a Gaussian rv with mean $\mu_1 + \mu_2$ and variance $(\sigma_1)^2 + (\sigma_2)^2$. A sampling method could take many model runs to establish this, despite it being an obvious result of the model structure. For this example, the GUM approach will be sufficient but if the input distributions were rectangular it would not. Analytical methods aim to include information from the model to enable calculations to be more efficient and accurate.

In most cases the methods are not a direct route to an analytical expression for the pdf. They can be a way to reach a form that still requires a sampling method to evaluate the pdf, but often the form that is reached is simpler to simulate than the original form due to the use of features of the original model.

The main drawback with analytical methods is their difficulty. In some cases they are difficult to understand, they are often tricky to implement, and in some cases they can only be applied to fairly simple problems. Nevertheless, they have still been applied to metrology problems successfully [32], and it is possible that their use may become more simple and widespread if they are implemented in software packages.

Three types of method are described here: stochastic differential equations, emulators, and series expansion techniques. Stochastic differential equations aim to describe uncertainty by defining a calculus that is compatible with the evolution of uncertain quantities that are not differentiable by traditional means. Emulators are effectively approximation methods, replacing a computationally expensive computer model with a statistically equivalent model. Series expansion methods aim to derive analytical expressions for the statistics of the output quantities by rewriting the model in terms of summations of basis functions, with the hope that the sums will converge after sufficiently few terms for this to be computationally efficient.

Stochastic differential equations in particular are described at length for several reasons. The first reason is that they are likely to be unfamiliar to most readers, whereas it is likely that many metrologists have encountered Monte Carlo simulation before. A second reason is that it is quite difficult to understand SDEs without a large body of background material, so this material has been included as an appendix in section 10. A third is that SDEs are not currently particularly widespread in metrology, so it is worth explaining them in depth in order to promote their usage.

5.1 Stochastic differential equations

A stochastic differential equation (SDE) can be thought of as an ordinary differential equation with a forcing term that is described by a random variable. If the forcing is irregular, arising from some stochastic process such as from Gaussian white noise, then the solutions are not differentiable and new techniques are required (in fact a new calculus is required). In the physics literature SDEs are often called Langevin equations, particularly Hamiltonian equations with stochastic forcing, although the original Langevin equation is specifically

$$\frac{d\mathbf{v}}{dt} = -\lambda\mathbf{v} + \mathbf{f}(t)$$

where $\mathbf{f}(t)$ is a random forcing term.

One of the most useful applications of SDEs in a metrology context is the ability to calculate coverage intervals for quantities derived from noisy data [32]. This is discussed more in section 5.1.2 and case study 2. SDEs also occur in models of quantum systems, which are probabilistic by nature. They are also commonly used in financial applications to model currency prices, where a good estimate of variance is particularly important so that future option trades can be advantageously priced.

The mathematics that describes stochastic differential equations is explained more fully in section 10. It is not described in the main body of the text because it requires quite high-level mathematics and may not be of interest to the casual reader.

An ordinary differential equation $u' = f(t, u)$ can be solved to get an integral equation:

$$u(t) = u(t_0) + \int_{t_0}^t f(s, u(s)) ds.$$

This formulation is adapted for a system that is coupled to some random forcing. Symbolically the system is often written as a differential equation

$$dX(t) = a(t, X(t))dt + b(t, X(t))dW(t), \quad X(0) = X_0, \quad 0 \leq t \leq T, \quad (2)$$

where $b(t, X(t)) dW(t)$ is the random forcing. This is a notation convention, since the forcing term $b(t, X(t)) dW(t)$ is not in fact differentiable. The term $a(t, X(t))$ is called the drift, and $b(t, X(t))$ is called the diffusion.

The system (2) is more properly interpreted as

$$X(t) = X_0 + \int_0^t a(s, X(s))ds + \int_0^t b(s, X(s))dW(s), \quad 0 \leq t \leq T. \quad (3)$$

The second integral on the right-hand side is integration with respect to the noise. Throughout the following, it will be assumed that the noise is Brownian motion, discussed further in section 10.1.1, although case study 2 involves coloured noise that is not described by simple Brownian motion. Additionally, the stochastic process will be denoted as $X(t)$, although X_t is also commonly used. This choice avoids confusion with the notation for “derivative with respect to t ”. X is assumed to be a function of a single variable, although the case for multi-dimensional problems is considered in section 10.1.3.3.

If b is independent of X then the noise is called **additive noise** otherwise it is called **multiplicative**. From a metrology point of view the noise term is used to model noise in the physical system. Additive noise is noise that is external to the system, whereas multiplicative noise depends on the variable X and can be thought of as internal to the system.

For example, Kloeden and Platen [33] consider an example from radio astronomy [34] where additive noise is used to model measurement error. The signal from a star is represented as

$$\eta(t) = A \exp(i(B + X(t))) + r\xi(t),$$

where A is the amplitude and B is the signal phase. $X(t)$ is real-valued with mean zero and is used to model atmospheric turbulence as a stochastic process, $\zeta(t)$ is complex-valued Gaussian white noise (Gaussian white noise is explained in section 10.1.1) which models measurement errors, and r is a parameter.

An Ornstein-Uhlenbeck (OU) process is used to model the measurement error $X(t)$:

$$dX(t) = -\beta X(t)dt + \sigma\sqrt{2\beta}dW(t),$$

where W is Gaussian white noise. OU processes generate coloured noise. Methods of noise generation and process simulation for such processes are explained further in the second case study (see section 7.2).

Measurement data is then used to determine estimates of the values of A and B for a star. The determination process also requires knowledge of β , σ , and r . If these parameters are known, one way to find A and B is to consider them as time-independent processes that satisfy

$$dA(t) = 0, \quad dB(t) = 0,$$

and then use a non-linear filter to determine estimates of A and B from the observations. Filtering is explained further in section 5.1.1.

5.1.1 Calculating statistics and estimating parameters using SDEs

Consider the equation

$$dX = \alpha a(X(t))dt + dW(t),$$

and suppose that α is a parameter to be estimated and a is a (possibly nonlinear) function. If X has been measured during an interval $[0, T]$, then parameter estimation methods seek the value of α that maximises the **likelihood ratio**

$$L(\alpha, T) = \exp\left(\frac{\alpha^2}{2} \int_0^T a^2(X(t))dt - \alpha \int_0^T a(X(t))dX(t)\right).$$

This can be justified heuristically, and is explained more fully in Kloeden and Platen [35]. The maximum likelihood estimator from this expression is

$$\hat{\alpha}(T) = \frac{\int_0^T a(X(t))dX(t)}{\int_0^T a^2(X(t))dt}, \quad (4)$$

which is a function of T because its value will depend on the integration limits. It can be shown that this expression will converge to the true value of α as $T \rightarrow \infty$ and an expression for the variance of the estimator exists.

An approximation to this expression can be calculated from discrete observations x_i , $i = 1, 2, \dots, N$ of the process,

$$\tilde{\alpha}_N = \frac{\sum_{i=1}^N a(x_{i-1})(x_i - x_{i-1})}{\sum_{i=1}^N a^2(x_{i-1})(t_i - t_{i-1})}, \quad (5)$$

and an example of this technique applied to real data is shown in figure 10. In this case, $a = X(t)$ and α converges to a value of -0.0342 s^{-1} after a time $T = 146\text{s}$ ($N = 664$). These data are measurements taken during the measurement process described in case study 2.

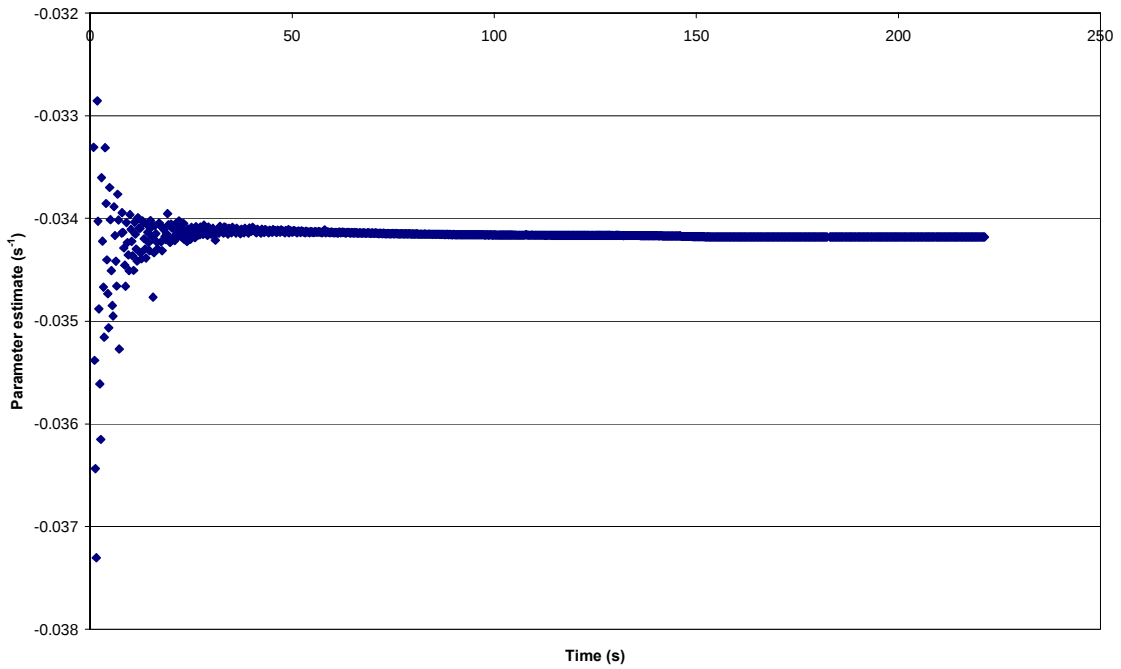


Figure 10: Convergence of a parameter estimate calculated using equation (5). The data are measurements gathered using the process described in case study 2.

Estimates of the mean and standard deviation of a random variable can be calculated using repeated realizations of the process, similar to sampling methods outlined in section 4. Then if estimates are calculated from n sample paths,

$$T_n = \frac{\hat{\mu}_n - \mu_n}{\sqrt{\hat{\sigma}_n^2/n}}$$

is distributed by the t-distribution with n degrees of freedom. This process is demonstrated in case study 2.

Under some circumstances, it is possible to obtain differential equations for moments of the random variable such as the mean and variance [31]. However, it should be noted that obtaining moments of order higher than 2 is usually very difficult, and so construction of anything more complicated than means and variances is usually too complicated to consider.

The use of stochastic differential equation techniques to derive parameters from noisy data is a form of filtering. The techniques usually involve maximum likelihood estimators and approximate likelihood estimators. Filtering is discussed in [35], and in

particular an example is given of Kalman filtering for a linear SDE of the form

$$dX(t) = AX(t)dt + BdW(t),$$

where A and B are constants. The filtering is like finding a conditional probability and is related to the Fokker-Planck equation.

Without going into detail, the approximate procedure for estimating the value of some parameter θ is to maximise the likelihood function of θ , which is given by the conditional probability product

$$\Lambda(x(t_1), x(t_2), \dots, x(t_N), \theta) = p(x(t_1) | \theta) \prod_{i=1}^{N-1} p(x(t_{i+1}) | x(t_i), \theta)$$

where $x(t_i)$ are the observed data points. Usually the logarithm of the likelihood is considered and the term $p(x(t_1) | \theta)$, which is asymptotically small, is neglected, so

$$L(x(t_1), x(t_2), \dots, x(t_N), \theta) = \sum_{i=1}^{N-1} \ln \{p(x(t_{i+1}) | x(t_i), \theta)\}.$$

This function is then differentiated and maximised by finding the stationary point.

This technique requires the conditional probability densities, which can be found using the Fokker-Planck equations as explained in section 10.1.3.2, or alternatively by integrating the SDE with a guess for the parameter and estimating the conditional density [36].

5.2 Emulators and response surface methodologies

Emulators and response surface methodologies try to replace a piece of computationally expensive computer code with a model that can be used as a surrogate for the code itself. Emulators [37] (also called “metamodels”) are often used for extremely large models such as climate simulations and environmental applications [38], which frequently feature the coupling of several different large models. A form of the technique is also used in geological industries (oil, mining etc.) where it is called kriging.

Emulators are developed by using regression techniques on model results. They aim to use Bayesian techniques to combine information about the model inputs with model results to construct a suitable model for the output quantities. One of the most common methods for carrying out this regression is the use of Gaussian process techniques [39]. This technique assumes that the output quantity $Y(\mathbf{X})$ is smooth and continuous, and uses these properties to infer the behaviour of Y in areas surrounding the data points $y_i = Y(\mathbf{x}_i)$. The value of Y is treated as an inexact quantity, and is written

$$Y(\mathbf{X}) = f(\mathbf{X}) + \sum_{i=1}^N w_i \phi_i(\mathbf{X}),$$

where $f(\mathbf{X})$ is a deterministic parametric function, the w_i are zero-mean weights with a multivariate Gaussian distribution, $\mathbf{w} \sim N(\mathbf{0}, \Sigma_w)$, and the ϕ_i are a set of basis functions. The form of f is chosen during the regression procedure, and the ϕ_i can be chosen to be any suitable set of basis functions.

Response surface methods are generally relatively straightforward to implement, but are

less likely to be accurate for strongly nonlinear problems. The methods fit a polynomial surface to model results using least-squares regression techniques. The order of the polynomial chosen generally depends on the number of model results available, but frequently low-order polynomials are regarded as sufficient.

These methods (particularly response surface methodology) are frequently unsuitable for extrapolation. This means that care must be taken when selecting the input quantity combinations that are used to provide the input data for the regression. It is best to use a sampling technique like experimental design or Latin Hypercube sampling (see section 4) to ensure that all areas of the input quantity domain are covered.

The main drawback with these methods is that it is always difficult to be certain that the emulator or response surface is simulating the model's behaviour sufficiently accurately. In particular, both emulators and response surfaces make assumptions about the smoothness of the behaviour of the quantities of interest. The methods are unsuitable for discontinuous or otherwise ill-behaved functions.

5.3 Series expansion methods

Series expansion methods can be used to approximate processes where it is believed that the uncertainties of the input quantities are sufficiently small to allow for expansion of either the output quantity or the operator as a sum in terms of the small quantities. They are usually applied to problems where the uncertain input quantities have led to the differential operator having an uncertain form, so for instance if the heat equation with forcing $\nabla \cdot (\lambda \nabla T) = f(\mathbf{x})$ has an uncertain spatially-varying λ , then the differential operator is no longer deterministic.

Suppose that the equation of interest can be written as $Lu = f$, where L is an uncertain linear differential operator, u is the output quantity, and f is the forcing (assumed in this case to be deterministic). Then if L has a deterministic part and a probabilistic part, the equation can be written as

$$(\Lambda + \Pi)u = f \quad (6)$$

where Λ is the deterministic part and Π is the probabilistic operator. Λ includes all the deterministic terms, so the probabilistic operator must have mean zero. As an example, consider the one-dimensional forced heat equation

$$\frac{\partial}{\partial x} \left(\lambda(x) \frac{\partial T}{\partial x} \right) = f(x)$$

where λ is an uncertain quantity, f is a deterministic function, and T is the temperature. Writing $\lambda(x) = \lambda_0(x) + \varepsilon$, and assuming that ε a random variable of mean zero that is independent of position, gives

$$\frac{\partial}{\partial x} \left((\lambda_0(x) + \varepsilon) \frac{\partial T}{\partial x} \right) = \frac{\partial \lambda_0}{\partial x} \frac{\partial T}{\partial x} + \lambda_0(x) \frac{\partial^2 T}{\partial x^2} + \varepsilon \frac{\partial^2 T}{\partial x^2}$$

so that

$$\Lambda = \frac{\partial \lambda_0}{\partial x} \frac{\partial}{\partial x} + \lambda_0(x) \frac{\partial^2}{\partial x^2}, \quad \Pi = \varepsilon \frac{\partial^2}{\partial x^2}$$

are the terms in the decomposition.

The simplest series expansion method is perturbation analysis. The operator Π and the uncertain quantity u are expanded in terms of the uncertainties of the input quantities. Suppose that the input quantities are $x_i = \mu_i + \alpha_i$, $i = 1, 2, \dots, n$, where the μ_i are the mean values and the α_i are the uncertainties. Then the expansions are of the form

$$\Pi(\mathbf{a}, \mathbf{r}) = \sum_{i=1}^n \alpha_i \frac{\partial \Pi}{\partial \alpha_i} + \sum_{j=1}^n \sum_{i=1}^n \frac{\alpha_j \alpha_i}{2} \frac{\partial^2 \Pi}{\partial \alpha_i \partial \alpha_j} + \dots$$

$$u(\boldsymbol{\mu}, \mathbf{a}, \mathbf{r}) = \bar{u} + \sum_{i=1}^n \alpha_i \frac{\partial u}{\partial \alpha_i} + \dots$$

where $\mathbf{r}$ is the position vector. This will be a good approximation if the uncertainties are much smaller than the mean values. This results in a series of equations that supply terms in the expansion of u . The first two terms are

$$\Lambda(\mathbf{x}, \boldsymbol{\mu}) \bar{u} = f$$

$$\Lambda(\mathbf{x}, \boldsymbol{\mu}) \frac{\partial u}{\partial \alpha_i} + \frac{\partial \Pi}{\partial \alpha_i} \bar{u} = 0.$$

These expressions give the first two terms in the approximation for u , and so an approximation to the pdf can be assembled. The larger the α_i , the more terms are required in the expansion. Some of the higher order terms contain expressions that cause the sum to diverge, which means that higher order approximations should not be used. The method may also be used on non-linear equations, but the resulting expansions will be considerably more complicated.

This is a form of Taylor expansion about the mean values, and clearly it will work best for small uncertainties. Generally a coefficient of variation of 10-15% is the maximum allowable, and the method works best if the input distributions are approximately Gaussian [40].

The expansion (6) can also be solved by application of a Neumann series expansion. The Neumann series expansion for the inverse of an operator states that

$$(I - M)^{-1} = I + M + M^2 + M^3 + \dots,$$

where M is a linear operator and I is the identity operator. If Λ is inverted, the problem can be written as

$$(I + \Lambda^{-1} \Pi) u = \Lambda^{-1} f$$

and so

$$u = (I + \Lambda^{-1} \Pi)^{-1} \Lambda^{-1} f$$

$$= (I - [\Lambda^{-1} \Pi] + [\Lambda^{-1} \Pi]^2 - [\Lambda^{-1} \Pi]^3 + \dots) \Lambda^{-1} f.$$

Note that this form removes the need to invert a probabilistic operator. The series will converge if

$$\|\Lambda^{-1}(\boldsymbol{\mu}, \mathbf{x}) \Pi(\boldsymbol{\mu}, \mathbf{x}, \mathbf{a})\| < 1.$$

Calculation of the sum is usually terminated after the first two terms since the higher order terms rapidly become complicated. Additionally, evaluation of the probabilistic

operator will require a sampling method to calculate the mean and variance of u . The Neumann expansion method can be extended and improved, as explained in Ghanem and Pol [41].

A more advanced series expansion method is stochastic finite element analysis (SFEA) [18, 41]. This method is fully explained in Ghanem and Pol [41], including worked examples and a method for deriving statistical moments from the results. The method expands the Karhunen-Loeve series, which is a series of the form

$$w(\mathbf{x}, \theta) = \sum_{n=0}^{\infty} \xi_n(\theta) \sqrt{\lambda_n} f_n(\mathbf{x})$$

where θ denotes a probabilistic quantity, the ξ_n are random variables that can be determined, and λ_n and f_n are the eigenvalues and eigenfunctions of the (known) covariance function of the input quantities. Examples of eigenvalues and eigenfunctions are given for common covariance functions in Ghanem and Pol [41]. This expansion is fine for the probabilistic operator but is not suitable for the unknown quantity of interest since its covariance function is unknown. The equivalent expansion for the unknown quantity takes the ξ_n and expands them as an expansion in terms of a set of polynomial functions of random variables with certain restrictions, called a Polynomial Chaos. The spatial parts of the expansions are then approximated using a finite element method, and results can be derived for the probabilistic behaviour.

There are two main drawbacks with these methods. The first is that in general, they require perturbations of the input variables from their means to be small so that the series converges after a couple of terms. If this is not the case, the computation of the terms in the expansion rapidly becomes intractable. The second drawback is that the calculation of terms can become very complicated, particularly for non-linear equations. The SFEA methods outlined here involve high-level mathematics and are not immediately accessible to FE users. It is unlikely that such methods are available in any commercial FE packages at present. Nevertheless, such methods are being applied to real problems [42] and so may become more widely applied in the future.

6 Methods developed for particular applications

The methods described in this section were originally developed for particular applications. In some cases it is not possible to apply the techniques to other areas, but some of the techniques can be applied more generally, particularly if used in combination with a more general method such as MCS.

6.1 Statistical Energy Analysis

Statistical Energy Analysis (SEA) was developed in order to investigate the uncertainty in results from vibrational calculations [33]. Vibration calculations are common in the testing of large structures such as buildings, bridges, and aeroplanes, and in noise assessment for cars and similar products. For low-frequency vibration, finite element analysis can be used to calculate the deformation and energy distribution of the structure. For high-frequency vibration it is necessary to use a different method since FE calculations at high frequency require dense meshes and can become extremely computationally expensive. SEA does not require a mesh and so provided knowledge of the deformation behaviour is not required, SEA is suitable for calculations at high frequency.

SEA does not predict the deformation of a vibrating structure in detail. Instead, built-up systems are tackled by dividing them into subsystems and applying a power balance equation to each of them. If E_i is the vibrational energy in the i^{th} subsystem,

$$P_{in,i} = \omega \eta_i E_i + \sum_{k=1}^N \omega \eta_{i,k} n_k \left(\frac{E_i}{n_i} - \frac{E_k}{n_k} \right),$$

where P_{in} is the power in, ω is the average frequency, n_i is the modal density of the subsystem, N is the number of subsystems, and the η are loss factors (loss coupling factors if two subscripts are present). If the power, loss factors, and modal density are all supplied, then the equations above form a linear system and the energy in each subsystem can be calculated. If the power and the loss coefficients in the equations are regarded as randomly distributed, expressions for the statistics of the energy can be calculated and so the uncertainty of the energy distribution can be estimated.

For mid-range frequencies, SEA could be coupled with finite element analysis (FEA). In general, a structure with mid-range fundamental frequencies will consist of stiff components and flexible components joined together. If the stiff components are modelled using SEA and the flexible components are modelled using FEA, the two methods could be coupled to produce a method that is efficient for both types of structure. A different formulation of this coupling also makes it possible for the probabilistic process to be modelled with SEA and the deterministic part to be modelled with FEA. This is an area of ongoing research since at present it is not clear what form the coupling should take for best results.

The main drawback with the method is its limited applicability. It can only be applied to vibrational problems, and is best suited to vibration at high-frequency.

6.2 Most Probable Point methods

Most Probable Point (MPP) methods were originally developed for reliability testing and failure testing. In their original form, they aim to calculate the probability of failure of a structure when the failure state is defined by some condition. Two other common techniques related to MPP methods [22] are the first-order reliability method (FORM),

and the second-order reliability method (SORM). They are also called limit-state methods. These are not really suitable for uncertainty analysis but are common in reliability and sensitivity testing, particularly for structural analysis.

MPP methods evaluate the probability of failure where failure is defined by some limit function $g(\mathbf{X})$, so that $g(\mathbf{X}) \leq 0$ is failure and $g(\mathbf{X}) > 0$ is safe. The “most probable point” in the name is the point at which $g(\mathbf{X}) = 0$ that is most likely to occur. Usually g is defined as $g(\mathbf{X}) = f(\mathbf{X}) - f_i$, so if the probability that $f(\mathbf{X}) \leq f_i$ for a series of values is calculated, the distribution function can be constructed.

MPP methods attempt to calculate the failure probability,

$$p_f = \int_{\Omega} \dots \int_{\Omega} f_X(\mathbf{X}) dV,$$

where f_X is the joint pdf of the input quantities, and Ω is the region where $g(\mathbf{X}) \leq 0$. This means there are two problems in implementing MPP methods: finding the most probable point and the region Ω , and calculating the integral.

There are two common approaches for finding the most probable point. The most common method is to use an optimisation method to identify the correct point. This has the advantage that it will probably produce derivative estimates as it proceeds, and these can be used to calculate approximate sensitivity coefficients. Alternatively, an approximate response surface model can be used (as described in section 5.2), whereby $f(\mathbf{X})$ is approximated by a linear combination of the X_i . This form is easier to invert to find Ω . The distribution function calculated using a region Ω given by a linear approximation is called the mean value first order method. The first estimate of the distribution function can be improved by substituting the calculated values of $\mathbf{X}$ into the full form of $f(\mathbf{X})$ to get an improved estimate of the value of g there. This improvement can be repeated until convergence to a region Ω is obtained. The method involving repetition is usually known as the advanced mean value (AMV) method.

The problem of calculating the integral is often addressed using Monte Carlo methods for integration. As was mentioned in section 4, these methods are often used for complicated multi-dimensional integrals, and details of their implementation can be found in *Numerical Recipes* [27].

There are several drawbacks with the method. The most obvious is that the methods are developed to test one level of probability at a time, so that calculating a reasonably accurate pdf will take a large number of runs, and each run may require many model evaluations to find the MPP and evaluate the integral. This large number of runs may make a sampling method (see section 4) preferable.

Another problem is that in their simplest form, these methods assume a unique point at which $g(\mathbf{X}) = 0$, so that an optimisation algorithm can be used to find the (by assumption unique) MPP. There may be several equally likely points at which $g(\mathbf{X}) = 0$, which will mean the algorithm may miss one out unless it is run repeatedly with different starting points. This is an important point to consider if the response function is likely to have several peaks.

7 Case studies

The case studies in the following section have been drawn from different areas of metrology and have been chosen because they illustrate the use of some of the methods described in this report. In some cases the uncertainty calculation shown is only part of the overall uncertainty budget. In others, the case study is a simplified first step towards addressing a more complex form of the problem.

In each case, the studies describe:

- The physical problem
- The mathematical formulation of the problem
- The sources of uncertainty in the problem
- The uncertainty calculation method and reasons for the choice
- The results of the calculation
- Conclusions and possible extensions

The authors would like to thank the following people for their assistance with the problems described in the case studies: Hilmi Kurt-Elli (Rolls-Royce), Simon Duane (Centre for Acoustic and Ionising Radiation, NPL), John Redgrove and Bufa Zhang (both Centre for Basic, Thermal and Length Metrology, NPL).

7.1 Case Study 1: Strain measurement in blades

7.1.1 The physical problem

The performance of a rotor blade is measured by putting it into a test engine and increasing the engine speed over time so that the forces acting on the blade vary. The blade has strain gauges at various points on its surface, which collect data during the test.

The strains are caused by a combination of loading due to centripetal acceleration and vibration loading as the fundamental modes of the blade are excited. The strain gauge data can be analysed to investigate to what degree each of the fundamental modes are being excited at different times and different engine speeds.

The fundamental modes of the blade are modelled using finite element analysis and the deformation and stress distribution within the blade are calculated for each of the modes. Using the degree to which each mode is excited, as calculated from the strain gauge data, an overall stress distribution can be calculated from the combination of the stresses in different modes, and a peak stress level can be found.

This peak stress level is then compared with fatigue data to ensure that the peak stress is less than the acceptable limit for operational use. The problem is to calculate an uncertainty estimate for the final comparison between the fatigue data and the peak stress in the blade.

7.1.2 The mathematical model

The mathematical model of the quantity of interest is

$$f = \sigma_{\text{fatigue}} - \sigma_{\text{peak}}, \quad \text{where}$$

$$\sigma_{\text{peak}} = \max \left\{ \sum_{i=1}^N a^i(\varepsilon_{m,j}^i, \varepsilon_{c,j}^i) \sigma_c^i(\mathbf{x}) : \mathbf{x} \in \Omega \right\}$$

where σ_{fatigue} is the stress taken from the fatigue data, σ_{peak} is the peak stress within the blade calculated from the FE model, σ_c is a calculated stress, $\varepsilon_{c,j}$ is the strain calculated at the j^{th} strain gauge, $\varepsilon_{m,j}$ is the strain measured by the j^{th} strain gauge, a^i are the “degree of excitation” coefficients (which are functions of the measured and calculated strains), $\mathbf{x}$ denotes spatial position, Ω represents the rotor blade, and the superscript i denotes the i^{th} natural frequency, of which there are N .

7.1.3 Sources of uncertainty

The main sources of uncertainty within the problem are:

- Uncertainty in the strain gauge data and its interpretation, caused by
 - calibration problems,
 - noise and uncertainty of the measurement,
 - difficulties in identifying which fundamental frequency is which, and
 - misplacement and misalignment of the strain gauges.
- Uncertainties in the FE modelling caused by
 - the difference between the model and reality,

- the unknown accuracy of the FE model (caused by discretisation errors, numerical errors, etc.), and
- misalignment of anisotropic materials leading to errors in material property definitions.
- Uncertainties in the material fatigue data, including
 - uncertainty of experimental measurements and
 - uncertainty from interpolation and extrapolation of data at lower temperatures to obtain data at the operating temperature of the test engine.

These sources of uncertainty will lead to σ_{fatigue} , σ_c , $\varepsilon_{m,j}$, and $\varepsilon_{c,j}$ being uncertain quantities. Of these, the continuous model produces σ_c , and $\varepsilon_{c,j}$ so these quantities will be investigated in this study.

7.1.4 Initial study

For this study, a number of simplifications are made in order to address the key issues within a reasonable amount of time. However, the methods used are sufficiently generic that they can be applied to more complex model formulations.

The FE model geometry is simplified to a flat plate, dimensions 0.1 m by 0.5 m by 0.005 m, and boundary conditions are to clamp one of the short edges to remove all rigid body modes. It is assumed that the blade is measured using three vertically-oriented strain gauges, placed at (0.025, 0.025), (0.085, 0.255) and (0.025, 0.475) (see figure 11 for illustration), and that five fundamental frequencies are sufficient to capture the blade's excitation behaviour. Finally it is assumed that the degree of excitation coefficients are given by

$$a^i = \frac{1}{3} \sum_{j=1}^3 \frac{\varepsilon_{m,j}^i}{\varepsilon_{c,j}^i}$$

where the subscript j indicates the strain gauge.

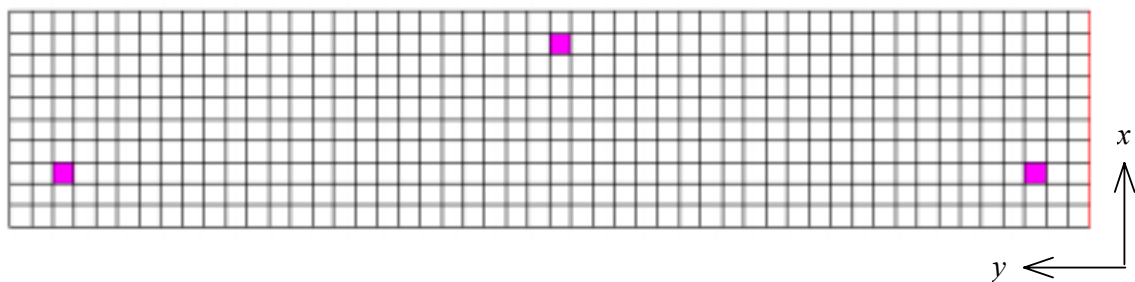


Figure 11: Mesh used for finite element analysis of the simplified blade. Pink squares show the locations of strain gauges, and the right-hand edge marked in red is clamped. Note the orientation of the axes.

The initial study is focussed on part of the first source of uncertainty mentioned in section 7.1.3: the misalignment and misplacement of the strain gauges. The misplacement and misalignment are characterised by three parameters: the offset Δx in position in the x direction, the offset Δy in position in the y direction, and the misalignment angle $\Delta\theta$. It is assumed that each of these quantities had a triangular distribution with mean 0, that they are independently distributed, and that the endpoints

of the distributions are ± 0.1 mm for Δx and Δy , and $\pm 5^\circ$ for $\Delta\theta$.

The strain gauges are modelled by considering the displacements of two points 0.01 mm apart and calculating the change in their separation relative to their original separation. This is assumed to be the strain that the strain gauge registers.

7.1.5 Choice of calculation method

Since the full-scale model is unlikely to be able to determine the strain without repeated runs, it is important to try and identify a process whereby good estimates of statistical quantities can be found using as few model runs as possible. An MCS run with a large number of trials is run, and the results of this simulation are treated as a reference solution for comparison purposes. Four different approaches are used to evaluate the uncertainty caused by the three parameters: an extreme value test, an orthogonal array of tests, tests using Latin Hypercube sampling, and a smaller Monte Carlo simulation. The latter two methods are run using a range of numbers of trials in order to gauge the behavioural trends as the number of trials varies.

7.1.6 Results of the calculations

Throughout the following, unless otherwise stated, the results shown are for a single frequency and a single strain gauge. The statements made have been checked for every frequency and gauge, and apply to all of the gauges at each of the frequencies. The strain gauge used is the one at (0.025, 0.025), and the frequency is the lowest one (15.68 Hz)

Table 5 shows the means and variances calculated using varying numbers of trials for the MCS and LHS trials. These results show that the mean and standard deviation are more accurate for small values of N if the LHS method is used instead of the MCS method. This conclusion is reinforced by the results shown in figure 12. Figure 12 shows the distribution functions calculated using 50 MCS trials, 50 LHS trials, and 10 000 LHS trials. It is assumed that the distribution function from 10 000 LHS trials is sufficiently close to the “true” distribution function that it can be used as a reference solution. The results from the 50 LHS trials are clearly closer to the reference solution than the results of the 50 MCS trials. From this table, it seems reasonable to take the mean value to be 2.33 and the standard deviation to be 0.72.

N	MCS		LHS	
	Mean (%)	Standard dev (%)	Mean (%)	Standard dev (%)
50	1.978	0.772	2.329	0.723
100	2.314	0.770	2.328	0.728
500	2.320	0.722	2.328	0.721
1000	2.333	0.718	2.328	0.721
5000	2.332	0.722	2.328	0.720
10 000	2.329	0.722	2.328	0.721

Table 5: Means and standard deviations of trials run using the MCS and LHS methods, with different numbers of trials N . All results reported above are strains and so are reported as percentages.

Figure 13 shows the difference between the distribution functions calculated from MCS

trials and those calculated from LHS trials for $N = 100$, 1000 and $10\,000$. This figure illustrates two points. The first point is that the difference decreases as N increases, so that the two methods tend towards the same limit as N tends towards infinity. The second point is that for this problem, for large values of N the largest error occurs at the very high and very low cumulative probabilities. This is probably because the highest and lowest cumulative probabilities correspond to the tails of the input distributions, and the LHS method ensures coverage of the tails where MCS does not.

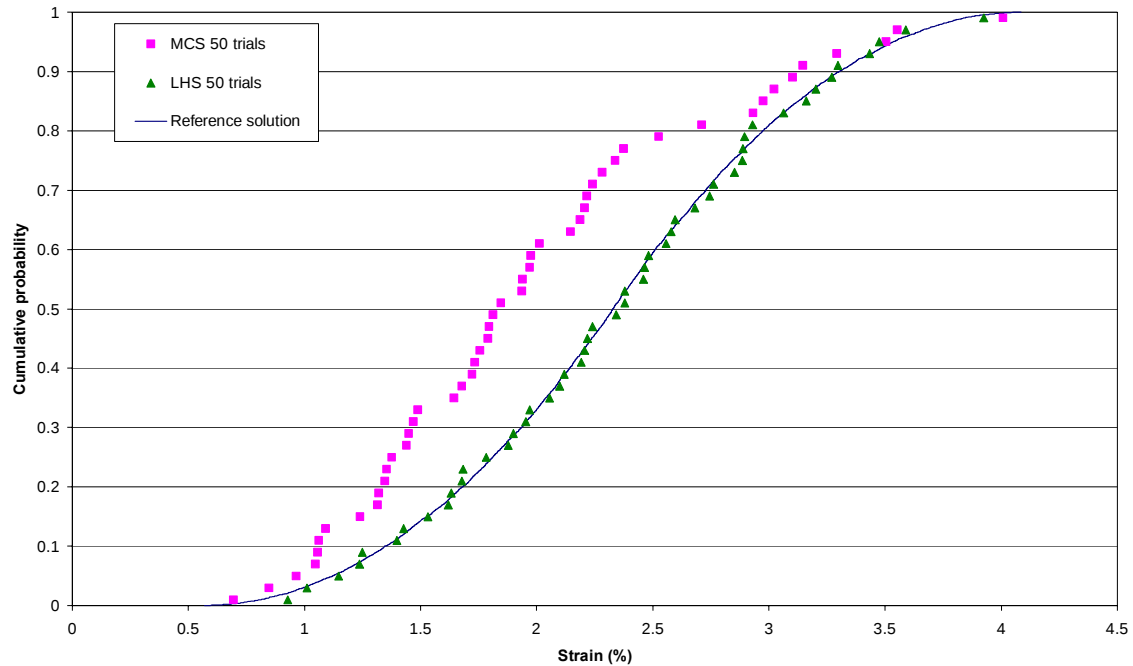


Figure 12: Distribution functions calculated from results of 50 MCS trials, 50 LHS trials, and 10 000 LHS trials

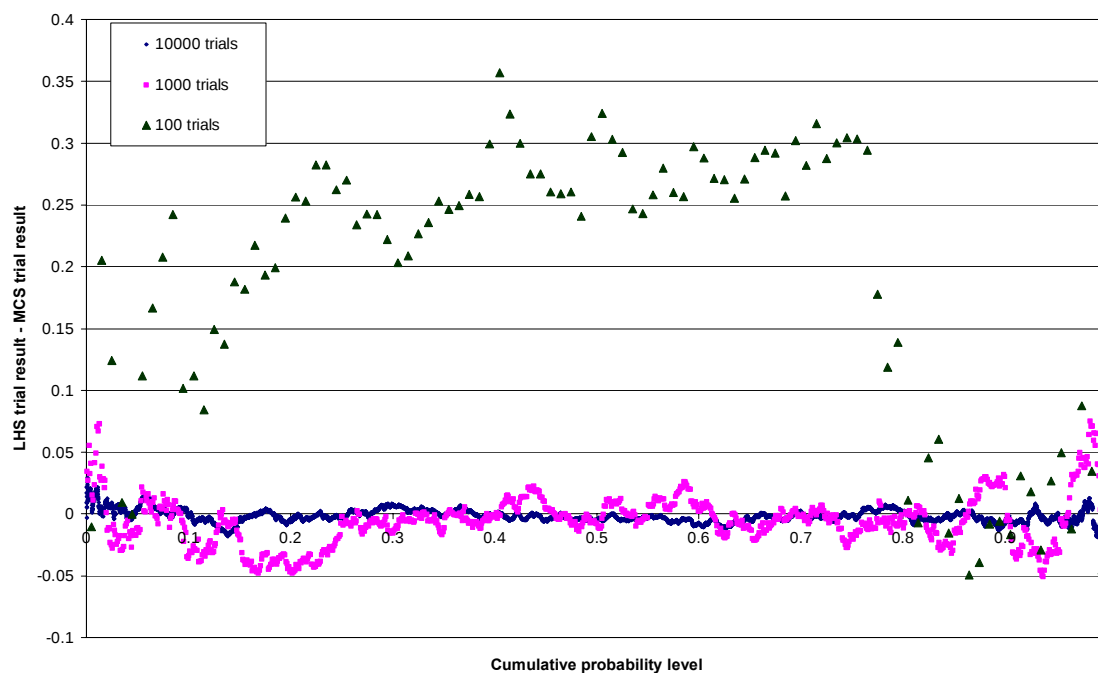


Figure 13: Difference between the LHS and MCS distribution functions plotted against cumulative probability for 10 000 trials, 1000 trials, and 100 trials.

The distribution function of the reference solution shown in figure 12 is a good fit to that of a triangular distribution. The pdf of the reference solution is shown in figure 14, and inspection shows that it resembles that of a triangular distribution. The figures in table 6 show the endpoints of the triangular distribution calculated as $\mu_Y \pm \sigma_Y/6$ from the various tests. The table also includes the most extreme values calculated from the DOE trials. The difference between the DOE trial values and the MCS and LHS values is probably due to the distribution not being exactly triangular at the tails.

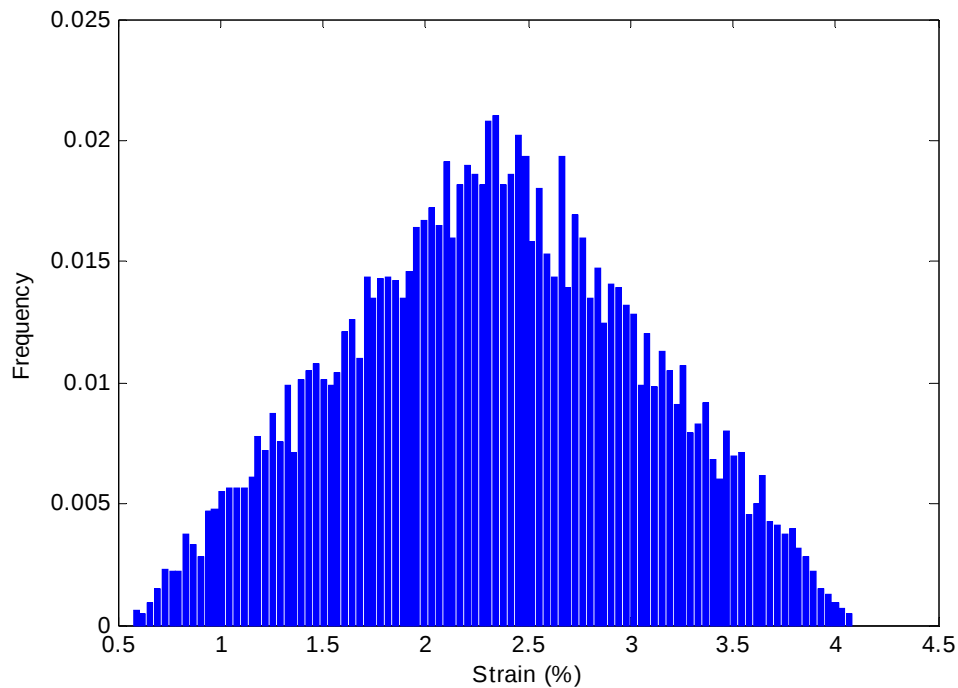


Figure 14: The pdf of the reference solution.

N	MCS x^- (%)	LHS x^- (%)	MCS x^+ (%)	LHS x^+ (%)
50	0.0880	0.5588	3.8686	4.0928
100	0.2534	0.5453	4.0247	4.0918
500	0.5510	0.5606	4.0894	4.0947
1000	0.5756	0.5605	4.0910	4.0946
5000	0.5628	0.5634	4.1012	4.1113
10 000	0.5598	0.5624	4.0987	4.0928
DOE: 9 trials	0.479		4.141	

Table 6: Extreme values for the triangular distribution with the same mean and standard deviation as the various MCS and LHS trials, compared to the extreme values calculated during the DOE trials.

7.1.7 Combining the uncertainties from different sources

As was stated in section 7.1.2, the quantity of interest is defined as

$$f = \sigma_{\text{fatigue}} - \sigma_{\text{peak}}, \quad \text{where}$$

$$\sigma_{\text{peak}} = \max \left\{ \sum_{i=1}^N a^i(\varepsilon_{m,j}^i, \varepsilon_{c,j}^i) \sigma_c^i(\mathbf{x}) : \mathbf{x} \in \Omega \right\}, \quad \text{and}$$

$$a^i(\varepsilon_{m,j}^i, \varepsilon_{c,j}^i) = \frac{1}{M} \sum_{j=1}^M \frac{\varepsilon_{m,j}^i}{\varepsilon_{c,j}^i},$$

where M is the number of strain gauges. Since this is essentially a discrete model, the GUM approach is used for evaluation of the uncertainty of f . The input uncertainties come from the following $2MN + N + 1$ quantities:

$$\sigma_{\text{fatigue}}, \left\{ [\varepsilon_{m,j}^i, \varepsilon_{c,j}^i, j = 1, \dots, M], \sigma_c^i, i = 1, \dots, N \right\}$$

and so the full expression for the uncertainty of f will be

$$u_f^2 = u_{\text{fatigue}}^2 + \max \left\{ \sum_{i=1}^N \left\{ \frac{(u_c^i)^2}{M^2} \left(\sum_{j=1}^M \frac{\varepsilon_{m,j}^i}{\varepsilon_{c,j}^i} \right)^2 + \left(\frac{\sigma_c^i}{M} \right)^2 \sum_{j=1}^M \left(\frac{\varepsilon_{m,j}^i}{\varepsilon_{c,j}^i} \right)^2 \left[\frac{(u_{m,j}^i)^2}{(\varepsilon_{m,j}^i)^2} + \frac{(u_{c,j}^i)^2}{(\varepsilon_{c,j}^i)^2} \right] \right\} : \mathbf{x} \in \Omega \right\},$$

where u is the uncertainty of the quantity with the same subscripts and superscripts (so that u_{fatigue} is the uncertainty of σ_{fatigue} , etc.).

At present, no data is available for the standard uncertainty of the measured strains, the calculated stresses, and the fatigue data. It is assumed that

- the measured strains have an uncertainty of 0.01% strain for all the strain gauges over all the frequencies,
- the standard uncertainty of the calculated stress is constant over $\mathbf{x}$ and has a value of 0.01 % of the mean value, and
- the fatigue data has an uncertainty of 5% of the mean value.

These assumptions mean that the only part of the sum that needs to be maximised is the calculated stress term, which simplifies the procedure.

The means and uncertainties for the calculated strains, as calculated from the FE results, are shown in table 7. A hypothetical set of strain measurements was created by forcing the free end of the model at a frequency of 150 Hz in the out-of-plane direction and calculating the resulting strain using FE. The mean values obtained for the different gauges and frequencies are shown in table 8.

Frequency (Hz)	Strain Gauge 1		Strain Gauge 2		Strain Gauge 3	
	Mean (%)	S. d. (%)	Mean (%)	S. d. (%)	Mean (%)	S. d. (%)
15.68	2.328	0.721	1.475	5.173	0.760	6.082
98.14	3.476	1.190	-3.342	0.730	-0.896	6.060
156.6	13.442	4.827	6.381	7.236	-0.612	0.964
275.4	-4.776	1.850	0.024	4.270	-1.387	5.987
301.4	0.160	1.76×10^{-3}	-0.0711	9.91×10^{-4}	8.84×10^{-3}	5.42×10^{-4}

Table 7: Means and standard deviations of strains calculated from the FE results for the first 5 fundamental frequencies. Strains are reported as percentages.

Frequency (Hz)	Strain Gauge 1		Strain Gauge 2		Strain Gauge 3	
	Measured value (%)		Measured value (%)		Measured value (%)	
15.68	0.270		0.108		2.25×10^{-3}	
98.14	0.308		-0.278		-0.0193	
156.6	0.0103		-1.34×10^{-3}		-1.01×10^{-3}	
275.4	1.499		0.0216		0.264	
301.4	0.000		0.0000		0.0000	

Table 8: Simulated measured strain values for the first 5 fundamental frequencies. Note that the chosen loading does not excite the 5th frequency. Strains are reported as percentages.

The calculation of the full uncertainty gave a very large value (of the order of 100 GPa), and by inspection the cause of this was the second strain gauge, which has a large coefficient of variation for the fourth frequency. For a real measurement, the strain gauge would not be placed in that location as it is a region of low strain and high sensitivity.

7.1.8 Conclusions

The Latin Hypercube sampling method has been shown to produce a better approximation to a pdf than the Monte Carlo simulation method when only a small number of trials are possible.

This Case Study is an example of a problem where the uncertainties calculated using a continuous model are just one contribution to the overall uncertainty. The method for determining the overall uncertainty of the quantity of interest is the standard GUM method, since the quantity of interest is described by a discrete model.

7.2 Case Study 2: Determination of the time constant of a calorimeter

7.2.1 The physical problem

A calorimeter consists of a central core of graphite separated from a larger body of material by an air gap, as shown in figure 15. The central core temperature is measured by an embedded thermistor that is monitored using a DC Wheatstone bridge.

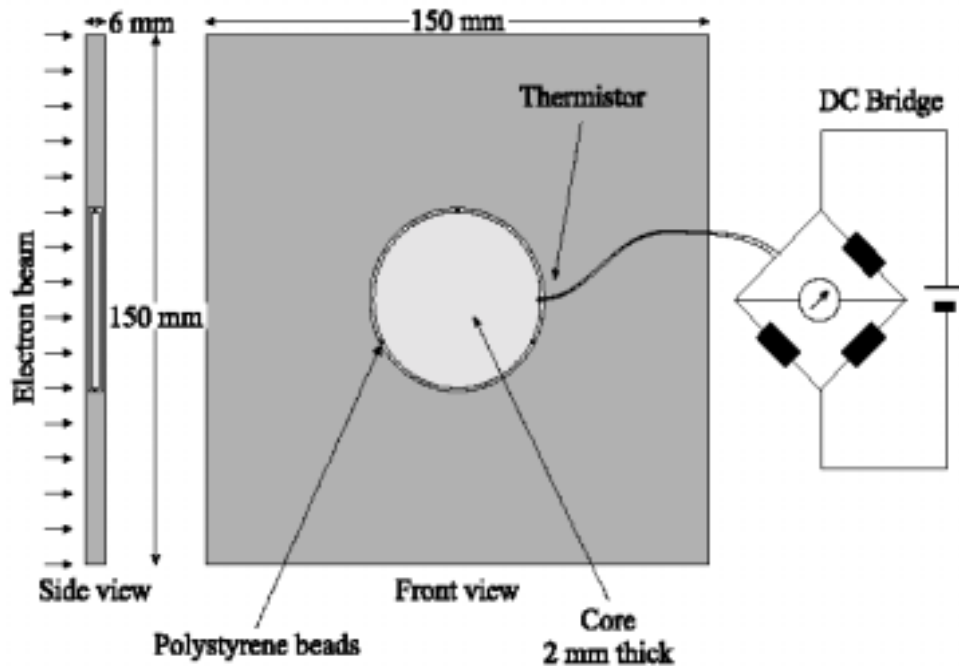


Figure 15: Calorimeter and measurement device.

When the calorimeter is irradiated with an electron beam, the temperature of the central core increases. The temperature changes, typically of the order of 1 mK, are such that electrical noise is not negligible. The measured signal changes occur over much shorter timescales than is possible for thermal effects, and these changes are due to noise generated in the thermistor.

The calorimeter is a portable device and it is important to ensure that it is working correctly when it is set up in a new location. One test is to heat the central core electrically and then let it cool, to extract the time constant of the core temperature exponential decay τ , and compare it to the value obtained in the laboratory. This comparison would depend on a reliable estimate of the uncertainties associated with both estimates of the time constant. The aim of the case study is to investigate how to obtain this estimate.

7.2.2 The mathematical model and choice of calculation method

The measured temperature $H(t)$ obeys

$$dH(t) = (\tau H(t) + \mu V(t))dt$$

where τ is a parameter controlling the exponential decay, μ is a parameter multiplying the noise term, and V is a random variable representing the noise.

A white noise term would satisfy the SDE

$$dV(t) = -\alpha V(t)dt + \eta dW(t), \quad (7)$$

where α and η are known parameters and W is a Wiener process (see section 10.1.1). Kloeden and Platen [35] show that the linear variable coefficient SDE

$$dX(t) = (a(t)X(t) + b(t))dt + c(t)dW(t)$$

has a solution that can be written

$$X(t) = \Phi_{t,0} \left(X(0) + \int_0^t \Phi_{s,0}^{-1} b(s) ds + \int_0^t \Phi_{s,0}^{-1} c(s) dW(s) \right) \text{ with } \Phi_{t,0} = \exp \left(\int_0^t a(s) ds \right).$$

This means that equation (7) will have solution

$$V(t) = e^{-\alpha t} \left(V(0) + \eta \int_0^t e^{\alpha s} dW(s) \right).$$

To evaluate this analytical solution, it would be necessary to calculate a numerical approximation to the stochastic integral using one of the methods mentioned in section 10.1.4, so numerical approximation is still necessary even though an analytical solution exists. Instead of using this solution, the governing system will be solved for V and H simultaneously.

The process (7) does not have the experimentally observed power spectrum. Analysis of measurement data has shown that the noise in this case is not Gaussian white noise (as was used throughout sections 5.1 and 10), but it is coloured noise. The techniques described in section 10 are still applicable to problems with coloured noise, because coloured noise can be produced by considering a system of SDEs with white noise rather than a single SDE. A better model for the process is given by the system of SDEs

$$dV_i = -\alpha_i V_i dt + dW_i, \quad i = 1, 2, \dots, N,$$

where the dW_i are independent Wiener processes. Taking

$$V = \sum_{i=1}^N a_i V_i$$

gives an improved approximation for a suitable choice of coefficients a_i and α_i . Analysis of measurement data shows that three terms may be enough to give a good approximation to the desired behaviour, provided that the coefficients are chosen carefully.

7.2.3 Sources of uncertainty

The main source of uncertainty that will be considered in this case is the noisy experimental data. The output quantity in this study is the parameter τ . An estimate of the uncertainty of τ can be calculated using the procedure outlined in section 5.1.1. The maximum likelihood estimator can be calculated from equation (5), and repeated evaluations of the estimator from different sets of data can give the statistics of its behaviour. This treatment of the estimator neglects the non-Gaussian nature of the noise, but it is thought that this does not have a significant effect.

7.2.4 Results of the calculations

Thirty data sets are analysed. Each set consists of observations H_i , $i = 1, 2, \dots, N$ at times t_i . For each set, the maximum likelihood estimator is calculated as

$$\hat{\tau} = \frac{\sum_{i=0}^{N-1} H_i (H_{i+1} - H_i)}{\sum_{i=0}^{N-1} H_i^2 (t_{i+1} - t_i)},$$

and the convergence of this estimate for a single typical data set with increasing N is shown in figure 10. It can be shown that if the estimator is calculated for M data sets, and the mean and standard deviation of these M calculations are μ and σ respectively, then $(\tau - \mu)/\sigma$ has a t-distribution with M degrees of freedom.

For the thirty data sets used, the mean value of the estimator was -0.00164 s^{-1} , and the standard deviation was $3.48 \times 10^{-4} \text{ s}^{-1}$. The distribution function and corresponding t-distribution are shown in figure 16, and they agree reasonably well. The root mean square difference between the distribution function probabilities and the t-distribution cumulative probabilities is 0.03.

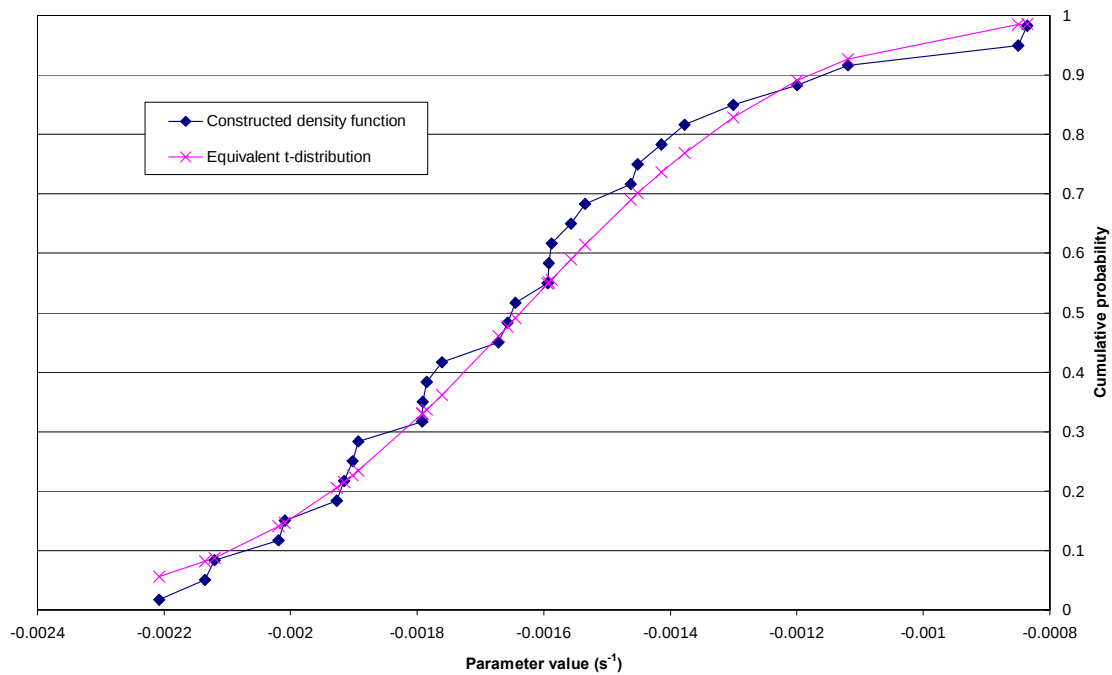


Figure 16: The density function of the parameter estimator calculated from thirty data sets, and the corresponding t-distribution.

7.2.5 Conclusions and extensions

Use of stochastic differential equation methods has led to an estimator for the time constant of a calorimeter, and has made evaluation of an associated uncertainty and distribution function possible.

Another application of the techniques described here is the use of SDEs for simulation purposes. The full system of SDEs could be solved using experimentally determined values for the parameters for many different noise sample paths, and the estimator could

be calculated for each simulation to give a distribution function like that shown in figure 16. Any differences between the simulated paths and real measurement data could help to identify any bias in the process for determining the time constant.

7.3 Case Study 3: Laser Flash Thermal Diffusivity

7.3.1 The physical problem

A cylindrical sample of material is held in a furnace in near-vacuum conditions so no heat is lost from the sample by convection. The sample is supported by three tiny pins, so no heat is lost to the surroundings by conduction. One of the circular faces (the “front face”) of the sample is exposed to a laser flash, and the temperature at the centre of the opposite face (the “rear face”) is measured as it rises due to the laser pulse and then falls due to radiative heat loss. This measured temperature data is processed to calculate a material property, the thermal diffusivity α . This is illustrated in figure 17.

The simplest form of the data-processing algorithm used to obtain α [43] makes various assumptions about the experimental set-up. The assumptions include: that there is no heat loss due to conduction or convection, that the material sample is isotropic, and uniform, and that the sample’s thermophysical properties are not affected by the small rise in temperature caused by the laser pulse. All of these assumptions are reasonable in practice. The algorithm also assumes that the laser pulse is instantaneous and uniform. These assumptions are less likely to be valid. In practice, the pulse is of finite duration, and the laser profile is not uniform, leading to uneven heating of the front face. The finite duration of the pulse is corrected by using an empirical expression in a more advanced form of the data processing algorithm.

The physical problem of interest is the evaluation of the contribution to the uncertainty caused by the non-uniformity and finite duration of the laser pulse.

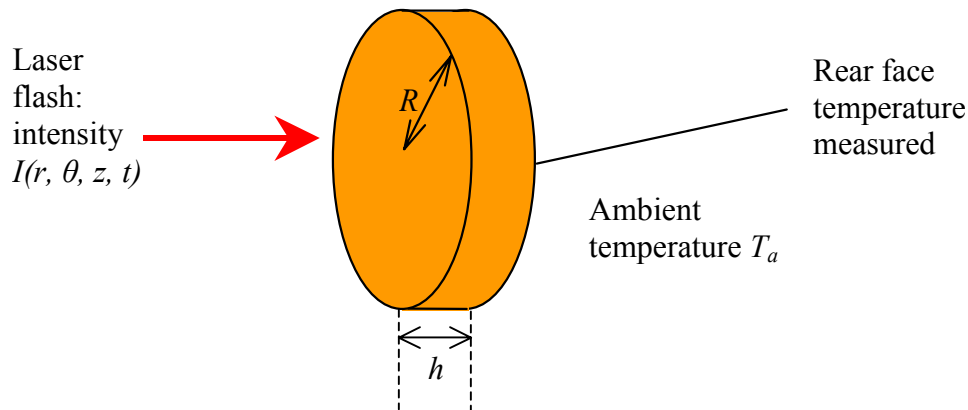


Figure 17: Illustration of the laser flash thermal diffusivity experiment.

7.3.2 The mathematical model

In its continuous form, the model of the experiment is

$$\frac{\partial}{\partial t}(\rho c_p T) = \nabla \cdot (\lambda \nabla T) + \delta(z) I(r, \vartheta, t) \quad 0 \leq r \leq R, 0 \leq z \leq h, 0 \leq \vartheta \leq 2\pi, 0 \leq t \leq t_0,$$

where λ is the thermal conductivity, c_p is the specific heat capacity, and ρ is the density, and the thermal diffusivity α is defined as $\alpha = \lambda / (\rho c_p)$. The boundary and initial conditions are

$$\begin{aligned}
T(r, \vartheta, z, 0) &= T_a, \\
\lambda \frac{\partial T}{\partial r} \Big|_{r=R} &= -\sigma \varepsilon_{\text{hem}} (T^4(R, \vartheta, z, t) - T_a^4), \\
\lambda \frac{\partial T}{\partial z} \Big|_{z=0} &= \sigma \varepsilon_{\text{hem}} (T^4(r, \vartheta, 0, t) - T_a^4), \\
\lambda \frac{\partial T}{\partial z} \Big|_{z=h} &= -\sigma \varepsilon_{\text{hem}} (T^4(r, \vartheta, h, t) - T_a^4).
\end{aligned}$$

where (r, θ, z) are cylindrical polar coordinates, ε_{hem} is the hemispherical emissivity, and σ_{SB} is the Stefan-Boltzmann constant. These equations describe transient heat transfer through a uniform cylinder with radiative heat loss from all faces and a heat source on the face $z = 0$.

This continuous model is then discretised and approximated using a fully three-dimensional finite difference model in cylindrical polar co-ordinates. The finite difference model makes no assumptions about laser uniformity; instead it accepts point-by-point definition of the laser intensity on the front face of the sample. It also allows for a non-instantaneous laser pulse by accepting an intensity vs. time curve. The model still assumes that the laser only affects the first layer of nodes directly. More details of the finite difference model can be found in a description of its validation process as a worked example in the report on model validation in continuous modelling [3].

The finite difference model produces an approximation to the transient temperature profile $T(r, \theta, z, t)$ which can then be processed using the same algorithm as that applied to experimental data, in order to determine an estimate of α .

7.3.3 Sources of uncertainty

The only source of uncertainty that is considered initially is the laser non-uniformity. The laser non-uniformity model is based on measurements made of two real laser profiles. The two profiles are thought to typify a “good” and “bad” laser profile in terms of uniformity. The chosen model is a two-parameter model of the form

$$I(x, y) = A + B \exp(-x^2/\sigma^2)$$

where x and y are Cartesian co-ordinates, B and σ are input quantities (B is dimensionless and σ is in mm), and

$$A = 0.01 - 0.1B\sigma \int_{-5/\sigma}^{5/\sigma} \exp(-x^2/\sigma^2) dx$$

so that $\int I(x, y) dx dy = 1.0$ over the area $-5 \text{ mm} \leq x, y \leq 5 \text{ mm}$.

The range of values that the parameters take in the measured profiles were $1.5 \leq \sigma \leq 5$ and $-0.036 \leq B \leq 0.01$. The parameters are taken to vary uniformly between these limits, since there are too few data points to construct any more reliable input distributions.

7.3.4 Choice of calculation method

The method chosen to address this problem is Monte Carlo simulation (as described in section 4.2). This method is chosen because the model typically takes a few minutes to run, only one result is required from each run, and it is straightforward to automate the entry of the input quantities. These factors mean that a large number of trials can be run automatically in a reasonably short time and results can be collected without generating

excessively large amounts of data. In order to simulate situations that occur in real experiments, two independent simulations are carried out. These sets of simulations use two different sets of material properties that are typical of the extreme values of properties seen in real materials: one set of properties for a material with a low thermal diffusivity (a “slow” material), and the other set for a metal with a high thermal diffusivity (a “fast” material). Each trial is run with 10 000 tests.

7.3.5 Results of the calculations

The distributions calculated for the fast and slow materials are shown in figures 18 and 19. The corresponding means, standard deviations, and 99% coverage intervals are given in table 9. The standard deviations of the estimates of the various quantities as described in section 4.2.2 have been calculated, and show that the estimates are acceptable if two figures are required for the uncertainty.

	Thermal Diffusivity	
	Fast material (m^2s^{-1})	Slow material (m^2s^{-1})
Mean value	1.13×10^{-4}	3.27×10^{-6}
50 percentile	1.14×10^{-4}	3.29×10^{-6}
Standard deviation	1.02×10^{-6}	2.71×10^{-8}
0.5 percentile	1.10×10^{-4}	3.21×10^{-6}
99.5 percentile	1.16×10^{-4}	3.35×10^{-6}
S_y	1.14×10^{-6}	3.29×10^{-8}
$S_{u(y)}$	2.09×10^{-8}	4.37×10^{-10}
$S_{y\text{low}}$	1.13×10^{-6}	3.25×10^{-8}
$S_{y\text{high}}$	1.16×10^{-6}	3.34×10^{-8}

Table 9: Means, standard deviations, 99% coverage interval end points, and standard deviations as explained in section 4.2.2, for the fast and slow materials, calculated from 10 000 MCS trials.

The shape of the distributions is not typical of any common distribution. This is reinforced by examination of a histogram approximation to the pdf, an example of which is shown in figure 20. The pdf vaguely resembles a union of three separate sub-regions, with a uniform distribution over each sub-region, but this is by no means clear. It is likely that this shape is due to some feature of the data processing algorithm. The data processing algorithm contains several “if ... then ... else ...” statements that are likely to produce distributions consisting of unions of distinct regions.

Another interesting point is illustrated by figures 21(a) and 21(b). These figures plot the thermal diffusivity against B and σ for the results from the slow material. They clearly show that the relationship between the diffusivity and the input quantities is non-linear and complicated. Attempts have been made to identify a function describing the relationship, but as yet none of the attempts have succeeded.

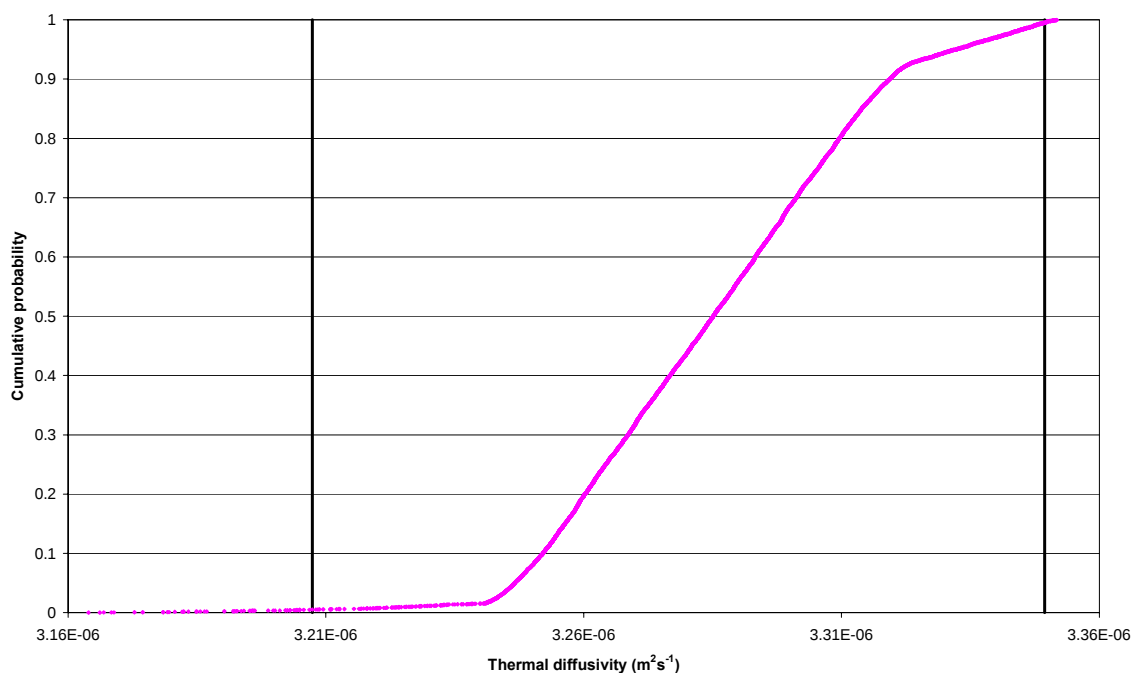


Figure 18: Calculated distribution function for the thermal diffusivity of a slow material. The solid black lines give a probabilistically symmetric 99% coverage interval.

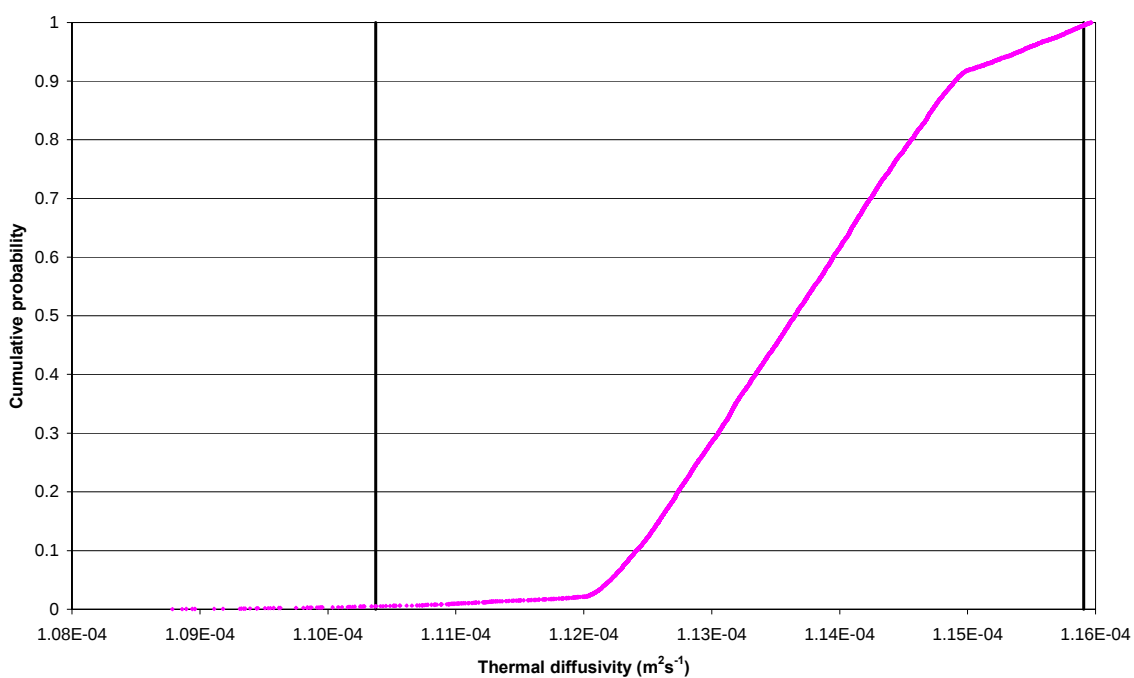


Figure 19: Calculated distribution function for the thermal diffusivity of a fast material. The solid black lines give a probabilistically symmetric 99% coverage interval. The similarity of this figure to figure 18 suggests that the shape is independent of material parameters.

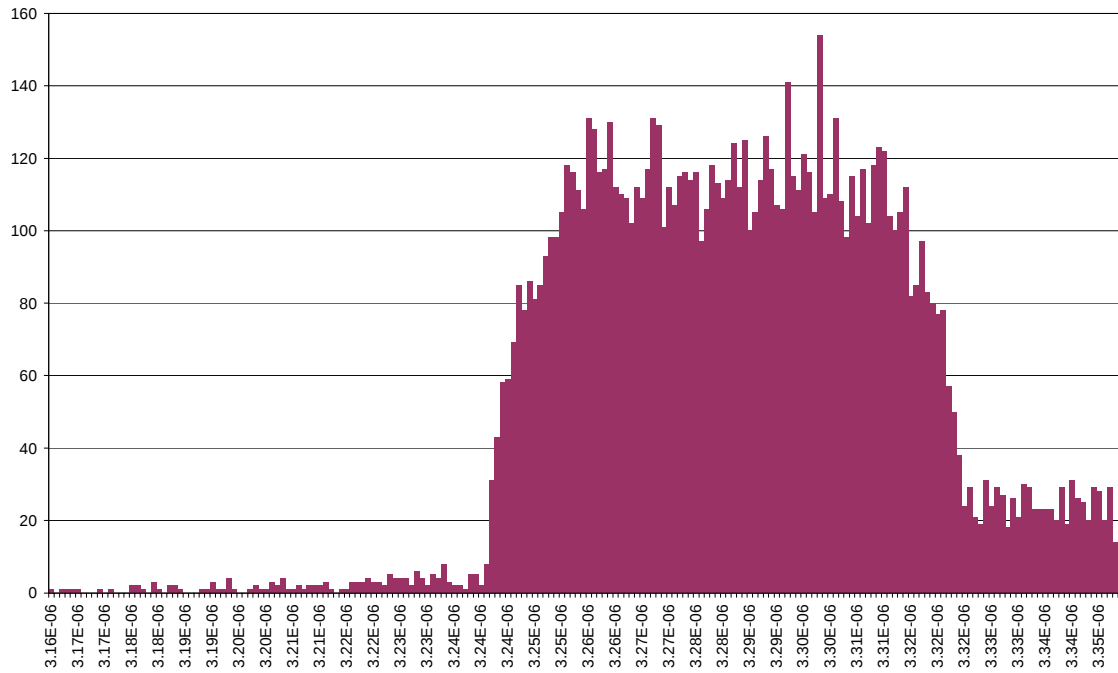


Figure 20: A histogram approximating the pdf of a slow material.

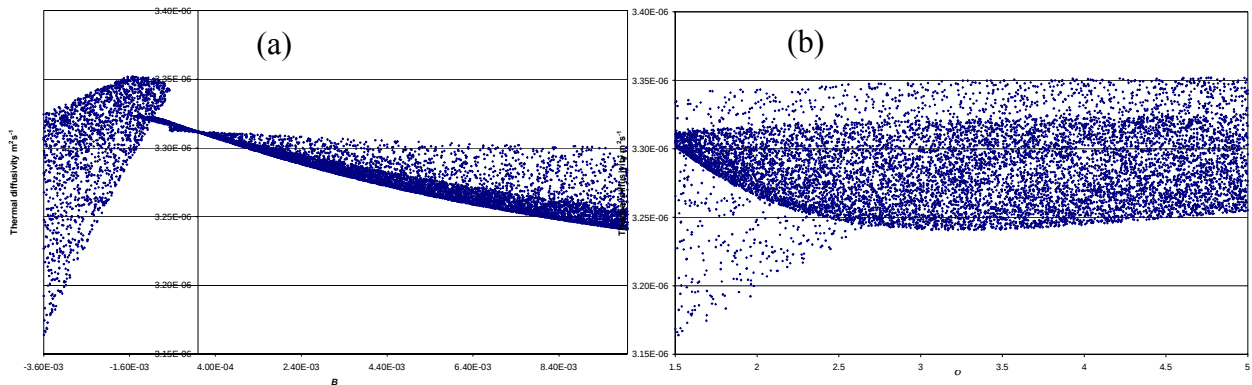


Figure 21: Dependence of thermal diffusivity (vertical axis) upon B (figure (a)) and σ (figure (b)) for a slow material. The dependence is clearly non-linear.

7.3.6 Conclusions and extensions

The distribution function for the thermal diffusivity does not resemble that of a normal distribution, and the dependence of the output quantity on the input quantities is clearly non-linear. This means that the Standard GUM [19] procedure would not have provided a good representation of the solution to the problem.

Monte Carlo simulation has produced a distribution function and statistical quantities. The results of a trial with 10 000 runs were adequate for a probabilistically symmetric 95% coverage interval but may not be sufficient for a probabilistically symmetric 99% coverage interval as the lower bound point has a large coverage interval. It may be possible to find a shorter unsymmetric 99% coverage interval whose endpoints have stabilised.

Further work under consideration includes the effects of a misaligned measurement spot on the rear face, and the use of the model for inverse problems such as determination of multiple material properties from experimental data.

8 Conclusions

Whilst techniques for uncertainty evaluation in continuous models are considerably less widely used than those used with discrete models, there are a number of methods that are potentially useful.

The GUM approach [19] is suitable for some problems, and for some models it may be a good way to find a rough estimate of the uncertainty from a simplified model before proceeding with a more complicated model that requires a different approach to uncertainty evaluation.

Sampling methods are readily applicable to continuous models. In particular, Monte Carlo Simulation (MCS) can be applied to any model and its use is becoming more widespread in uncertainty evaluation in discrete modelling [11, 12]. Many continuous models are computationally expensive, which makes the number of trials a method needs to determine a reliable uncertainty value an important issue. Alternative sampling methods such as Stratified sampling and Latin Hypercube sampling can provide good estimates of output quantity means and standard deviations using fewer trials than MCS.

Analytical techniques exist that use features of the continuous model such as its differential nature and the local continuity of solutions to evaluate uncertainties. Some of these methods require high-level mathematics, but they may become useful tools in the future if they are implemented in software packages.

Some methods have been developed for particular continuous modelling applications, and may not be applicable to problems outside of one area, but if a problem is of the right type, the most suitable may be application-specific.

The case studies have illustrated the application of the methods to real metrology problems and have compared different methods wherever possible.

Acknowledgements

This report has been produced as one of the deliverables of SS/M project 1.4 on Continuous Modelling in Metrology. The work was funded by the Department of Trade and Industry's National Measurement System Directorate.

The authors acknowledge the contributions made to the work by Peter Harris, (Centre for Mathematics and Scientific Computing, NPL), Hilmi Kurt-Elli (Rolls-Royce), Simon Duane (Centre for Acoustic and Ionising Radiation, NPL), John Redgrove and Bufa Zhang (both Centre for Basic, Thermal and Length Metrology, NPL).

9 References

- [1] "Model Validation in Continuous Modelling", T J Esward, G J Lord, and L Wright. NPL Report CMSC 29/03, 2003.
- [2] "A framework for assessing confidence in computational predictions", K M Hanson and F M Hemez, *Experimental Techniques* **25**, pp. 50-55, 2001.
- [3] "SLang – the structural language, a tool for computational stochastic structural analysis", C G Bucher, Y Schorling and W Wall, *Proc. 10th ASCE Eng. Mech. Conf.*, Boulder, pp1123-1126, 1995.
- [4] <http://www.uni-weimar.de/Bauing/ism/engl.Version/SLang/slang.Short.html>
- [5] ["A Toolkit For Uncertainty Quantification In Large Computational Engineering Models."](#), S F Wojtkiewicz Jr., M S Eldred, R V Field Jr., A Urbina, and J R Red-Horse, paper AIAA-2001-1455 in *Proceedings of the 42nd AIAA/ASME/ASCE/AHS/ASC Structures, Structural Dynamics, and Materials Conference*, Seattle, WA, April 16-19, 2001.
- [6] <http://endo.sandia.gov/DAKOTA/licensing/license.html>
- [7] "Capabilities and Applications of Probabilistic Methods in Finite Element Analysis," D S Riha, B H Thacker, D A Hall, T R Auel, and S D Pritchard, *Int. J. of Materials & Product Technology*, **16**(4/5), 2001.
- [8] <http://www.nessus.swri.org/publications.shtml>
- [9] "Adaptive Importance Sampling Method for the Stochastic Finite Element Analysis", S-H Kim and K-W Ra, *8th ASCE Specialty Conference on Probabilistic Mechanics and Structural Reliability*, 2000.
- [10] <http://www.profes.com/>
- [11] "Guide to the Expression of Uncertainty in Measurement, Supplement 1, Numerical Methods for the Propagation of Distributions", BIPM, IEC, IFCC, ISO, IUPAC, IUPAP, and OIML. To be published.
- [12] "Software Support for Metrology Best Practice Guide 6: Uncertainty and Statistical Modelling" M G Cox, M Dainton, and P M Harris, March 2001.
- [13] "A Comparison of Three Methods for Selecting Values of Input Variables in the Analysis of Output from a Computer Code", M D McKay, R J Beckman, and W J Conover, *Technometrics*, **21**(2), p239-245, 1979.
- [14] "Latin Hypercube Sampling for Stochastic Finite Element Analysis", A m J Olsson, G E Sandberg, *Journal of Engineering Mechanics*, **128**, p121-125, 2002
- [15] "Latin hypercube sampling as a tool in uncertainty analysis of computer models", M D McKay, *Proceedings of the 24th conference on Winter simulation*, p.557-564, December 13-16, 1992.
- [16] "Illustration of Sampling-Based Methods for Uncertainty and Sensitivity Analysis", J C Helton and F J David, *Risk Analysis* **22**(3), p591-647, 2001.
- [17] "An Evaluation of Methodologies for Uncertainty Analysis in Biological Waste Water Treatment", S G T Kops, P A Vanrolleghem, *Proceedings 9th Junior Scientist Workshop*, pp.6, April 11-14 1996.
- [18] "Stochastic Finite Element Method for Elasto-Plastic Body", M Anders and M

Hori, *Int. J. Numer. Meth. Engng*, **46**, pp.1897-1916, 1999.

[19] “Guide to the Expression of Uncertainty in Measurement”, BIPM, IEC, IFCC, ISO, IUPAC, IUPAP, and OIML. ISBN 92-67-10188-9, Second Edition.

[20] “A Methodology for Testing Spreadsheets and Other Packages Used in Metrology”, H Cook, M Cox, M P Dainton and P M Harris, NPL Report CISE 25/99, September 1999.

[21] T. J. Esward, L. Wright et al. Testing Continuous Modelling Software. NPL report in preparation, 2004.

[22] “Sensitivity Analysis”, ed. A Saltelli, K Chan, & E M Scott. Wiley, ISBN 0-471-99892-3, 2000.

[23] “Sensitivity Derivatives for Advanced CFD Algorithm and Viscous Modeling Parameters via Automatic Differentiation”, L Green, P Newman, and K Haigler, *Proc. 11th AIAA Computational Fluid Dynamics Conference, AIAA Paper 93-3321*, 1993.

[24] “Sensitivity Analysis for Atmospheric Chemistry Models via Automatic Differentiation”, A Sandu, G R Carmichael, and F A Potra, *Atmospheric Environment*, **31**(3):475-489, 1997.

[25] “Automatic Differentiation Techniques and their Application in Metrology”. R Boudjemaa, M G Cox, A B Forbes, and P M Harris, NPL Report CMSC 26/03, June 2003.

[26] “Software Support for Metrology Best Practice Guide 10: Discrete Model Validation”, R M Barker and A B Forbes, ISSN 1471-4124, March 2001. Available from <http://www.npl.co.uk/ssfm/download/abstracts.html#87>.

[27] “Numerical recipes in C/C++/F77/F90”, W Press, S A Teukolsky, W T Vetterling, and B Flannery, CUP.

[28] “Juran’s Quality Control Handbook 4th Edition”, J M Juran and F G Gryna (eds.), McGraw Hill, 1988.

[29] “Multivariate Analysis”, K V Mardia, J M Bibby and J T Kent, Elsevier Science & Technology Books, ISBN 0124712525, 1980.

[30] “A distribution-free approach to introducing rank correlation among input variables”, R L Iman and W J Conover, *Commun. Stat. Simul. Comput*, **11**(3), 311-334, 1982.

[31] “Matrix Computations”, G H Golub and C F van Loan, Johns Hopkins University Press, Third Edition, 1996.

[32] “Nonstationary Effects in Mode-Stirred Reverberation”, L Arnaut, Proceedings of EMC 2003, Zurich, 2003.

[33] <http://www.isvr.soton.ac.uk/DG/HighFrequencyVibrations.htm>

[34] “Estimation de paramètres dans les processus stochastiques en observation incomplète: Application à un problème de radio-astronomie”, F L Gland, PhD thesis, Univ. Paris IX, 1981.

[35] “Numerical Solution of Stochastic Differential Equations, no.23 in Applications of Mathematics”, P Kloeden and E Platen, Springer, Third edition, ISBN 3-540-54062-8, 1999.

- [36] “Density Estimation”, B Silverman, Chapman & Hall, London, UK, 1986.
- [37] <http://www.shef.ac.uk/~pas/research/bayes/model-uncertainty-research.html>
- [38] http://www.floodrisknet.org.uk/newsletters/2002-08/uncertainties_iam/view
- [39] "Learning and Inference in Graphical Models", ed. M I Jordan, Kluwer, 1998.
- [40] “Random Field Solutions Including Boundary Condition Uncertainty for the Steady-State Generalized Burgers Equation”, L Huyse and R W Walters, 2001, available from <http://techreports.larc.nasa.gov/icas/2001/icas-2001-35.pdf>
- [41] “Stochastic Finite Elements, A Spectral Approach”, R G Ghanem and P D Spanos, Dover Publications, ISBN 0-486-42818-4, 2003.
- [42] “Model Experiment and Numerical Simulation of Surface Earthquake Fault Induced by Lateral Strike Slip”, M Hori, M Anders, and H Gotoh, *Structural Eng./Earthquake Eng., Japanese Society of Civil Engineers*, **19**(2), p227-236, 2002.
- [43] “Flash method of determining thermal diffusivity, heat capacity, and thermal conductivity”. W J Parker, R J Jenkins, C P Butler and G L Abbott. *J. App. Phys.*, **32**(9), 1679-1684, September 1961.
- [44] “Strong convergence of an adaptive Euler-Maruyama scheme for stochastic differential equations”, H Lamba, J Mattingly, and A Stuart, Technical Report, University of Warwick, 2003.
- [45] “Software specifications for uncertainty evaluation and associated statistical analysis”. M G Cox, M P Dainton, and P M Harris. NPL Technical Report CMSC 10/01, National Physical Laboratory, Teddington UK, 2001.
- [46] “Exponential mean square stability of numerical solutions to stochastic differential equations”, D Higham, X Mao, and A Stuart, *LMS J. Comput. and Math.* **6**, pp.1-17, 2003.
- [47] “Stochastic differential equations and their numerical solution”, P Kloeden, Lecture notes from EPSRC Summer School, July 2000.
- [48] “An introduction to numerical methods for stochastic differential equations”, E Platen, ACTA Numerica, pp.197-246, 1999.
- [49] “Introduction to Stochastic Differential Equations”, T Gard, no.114 in Pure and Applied Mathematics, Marcel Dekker, New York, USA, 1988 ISBN 0-8247-7776-X.
- [50] “Stochastic Differential Equations”, B Oksendal, Springer, Fifth edition, 1998.
- [51] “An algorithmic introduction to numerical simulation of stochastic differential equations”, D Higham, *SIAM Review*, **43**, pp.525-546, 2001.
- [52] <http://www.maths.strath.ac.uk/~aas96106/algfiles>
- [53] “Maple for stochastic differential equations”, S Cyganowski, L Grüne, and P E Kloeden, in “Theory and Numerics of Differential Equations”, J Blowey, J Coleman, and A Craig, eds., Springer, 2000, pp.127-179, ISBN 3-540-41846-6.

10 Appendix I: Details of stochastic differential equation techniques

This section aims to give the background to the techniques described in brief in section 5.1. Much of this section consists of details of the theory underpinning stochastic differential equations. This explanation is necessary because stochastic differential equations depend on a redefinition of calculus that needs a full description if it is to be used.

10.1.1 Brownian Motion

Brownian motion is the most commonly used method of describing noise, and is probably the most commonly occurring stochastic process. Robert Brown, a Scottish botanist, observed that a small particle such as a seed in a liquid or gas moves about in a seemingly random way. The motion is explained by the constant collisions between the seed and the gas or liquid molecules. Physicists tried to measure the velocity of the particle and found that it varied wildly in magnitude and direction.

A scalar standard Brownian motion, also called a standard Wiener process over $[0, T]$ is a collection of random variables $W(t)$, which depend continuously on t in $[0, T]$, that satisfy

1. $W(0) = 0$ with probability 1,
2. For all $0 \leq s < t \leq T$, the random variable given by the increment $W(t) - W(s)$ is normally distributed with mean zero and variance $t - s$, or equivalently $W(t) - W(s) \sim N(0, t - s)$,
3. For all $0 \leq s < t < u < v \leq T$, the increments $W(t) - W(s)$ and $W(v) - W(u)$ are independent.

The sample paths of the Brownian motion are continuous but are not differentiable anywhere. From the definition of the Wiener process, the variance of $W(t)$ grows without bound as time increases whilst the mean remains zero. This means that as time increases the path must take larger positive and negative values and so oscillates about zero more and more. In fact,

$$\lim_{t \rightarrow \infty} \frac{W(t)}{t} = 0,$$

and sharper estimates of growth are possible.

It is simple to approximate a Wiener process. Dividing the interval $[0, T]$ up into N subintervals of equal step length Δt , approximating the behaviour of the Wiener process over each interval $0 = t_0 < t_1 < \dots < t_N = T$ and letting $W_j = W(t_j)$, gives that $W(0) = 0$,

$$W_{j+1} = W_j + dW_j, \quad j = 1, 2, \dots, N$$

where the dW_j are a sequence of independent random variables with $dW_j \sim N(0, \Delta t)$. Alternatively the process can be approximated by a scaled random walk, by constructing a step-wise continuous random walk $S_N(t)$ and taking equally probable steps of length $\pm\sqrt{\Delta t}$ at end of each of the subintervals, with the sign of the step being figuratively chosen by the toss of a coin. Figure 22 shows a typical simulated Wiener process and the average of 1000 such processes. Note that the average is close to 0, the mean value of the process.

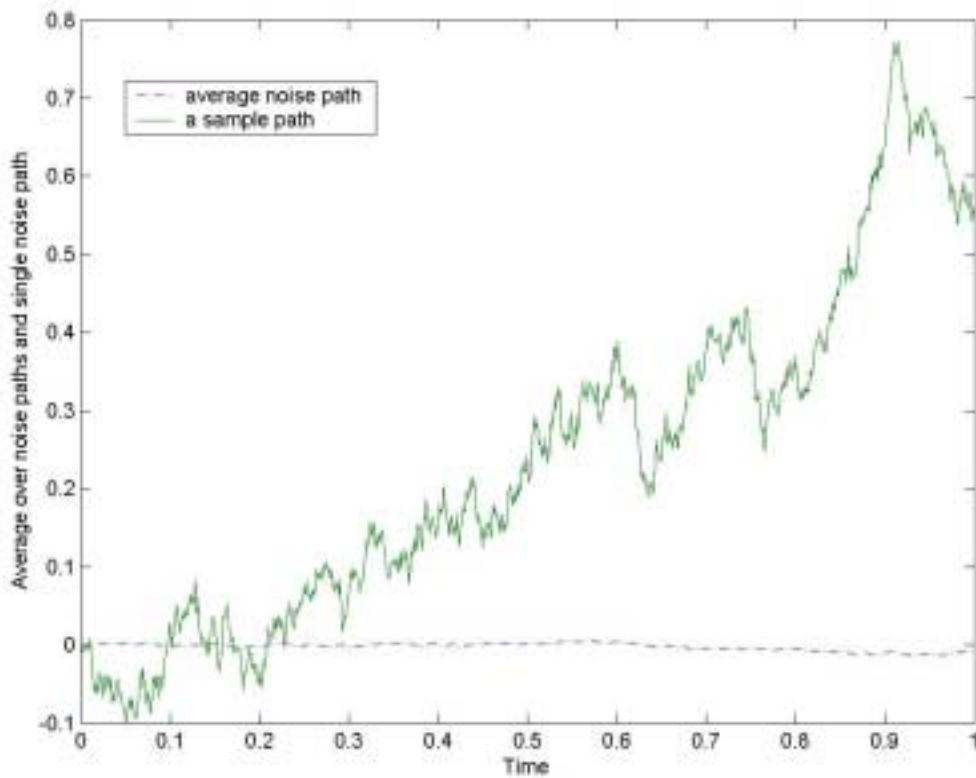


Figure 22: A sample path of a Wiener process (solid line) and the average of 1000 such sample paths (broken line). Parameter values were $T = 1.0$, $N = 1000$, $\Delta t = 0.001$

10.1.2 Stochastic calculus

The integral of a deterministic function g over an interval $[0, T]$ can be defined as the limit as $\Delta t \rightarrow 0$ of

$$\sum_{j=0}^{N-1} g(t_j^*) \Delta t \quad (8)$$

where $T = N\Delta t$ and $j\Delta t \leq t_j^* \leq (j+1)\Delta t$. Different choices of t_j^* will lead to the same value for the integral as $\Delta t \rightarrow 0$ by the definition of an integrable function as a function for which this limit exists and is unique. This uniqueness is not the case for stochastic integrals.

Two common definitions exist for

$$\int_0^T g(t) dW(t),$$

called the Itô calculus and the Stratonovich calculus. Both are based on sums of the form (8), and both have convergence in the mean square limit as $\Delta t \rightarrow 0$ (so that the convergence of the integral of g^2 can be shown). The **Itô** stochastic integral is defined as

$$\int_0^T g(t) dW(t) = \lim_{N \rightarrow \infty} \sum_{j=0}^{N-1} g(t_j) (W_{j+1} - W_j)$$

where *m.s. lim* denotes the mean square limit and $t_j = j\Delta t$. This uses the function value at the left-hand end of each subinterval. The **Stratonovich** stochastic interval uses the mid-point of each sub-interval for the function value, so that

$$\int_0^T g(t) \circ dW(t) = \lim_{N \rightarrow \infty} \sum_{j=0}^{N-1} g\left(\frac{t_j + t_{j+1}}{2}\right) (W_{j+1} - W_j).$$

Note that the symbol $\circ$ is used to denote that the integral should be interpreted using the Stratonovich formulation.

The Itô and Stratonovich integrals are different, even in the limit as $\Delta t \rightarrow 0$. It is in fact possible to develop a calculus based on the function value at any point in the subinterval $[t_j, t_{j+1}]$. Which form of integral is "correct" could be thought of as a modelling question due to the non-existence of a unique limit. The Itô calculus has some computational advantages. For instance, using the Itô calculus gives

$$E\left(\int_0^t g(X(s)) dW(s)\right) = 0, \quad E\left(\int_0^t g(X(s)) dW(s)\right)^2 = \int_0^t E(g(X(s)))^2 dW(s),$$

where X is a stochastic quantity. As an example of the difference between the two formulations, consider the difference between

$$\int_0^T W(t) dW(t) \text{ and } \int_0^T W(t) \circ dW(t).$$

Using the most general formula for the integrals,

$$\sum_{j=0}^{N-1} W(\theta t_j + (1-\theta)t_{j+1})(W_{j+1} - W_j),$$

where the Itô formulation has $\theta = 1$ and the Stratonovich formulation has $\theta = 1/2$. Since W is being sampled at discrete points, it is taken to behave linearly between those points, and so this sum is equal to

$$\begin{aligned} \sum_{j=0}^{N-1} \{\theta W_j + (1-\theta)W_{j+1}\}(W_{j+1} - W_j) &= \sum_{j=0}^{N-1} \left\{ \frac{1}{2}(W_j + W_{j+1}) - \left(\theta - \frac{1}{2}\right)(W_{j+1} - W_j) \right\} (W_{j+1} - W_j) \\ &= \frac{1}{2} \sum_{j=0}^{N-1} \{W_{j+1}^2 - W_j^2\} - \left(\theta - \frac{1}{2}\right) \sum_{j=0}^{N-1} (W_{j+1} - W_j)^2 \\ &= \frac{1}{2} (W_N^2 - W_0^2) - \left(\theta - \frac{1}{2}\right) \sum_{j=0}^{N-1} (W_{j+1} - W_j)^2. \end{aligned}$$

From the condition on the variance of W , the summation term in the final equation converges to Δt in the mean square sense, so that

$$\lim_{N \rightarrow \infty} E\left[\left(\sum_{j=1}^{N-1} (W_{j+1} - W_j)^2 - \sum_{j=1}^{N-1} \Delta t\right)^2\right] = 0.$$

Then, in the mean square sense, the integrals converge to

$$\frac{1}{2}(W(T)^2 - W(0)^2) - \left(\theta - \frac{1}{2}\right) \sum_{j=0}^{N-1} \Delta t = \frac{1}{2}W(T)^2 - \left(\theta - \frac{1}{2}\right)T.$$

This means that the solution using the Itô calculus will be $\frac{1}{2}(W(T)^2 - T)$ and the solution using the Stratonovich calculus will be $\frac{1}{2}W(T)^2$. This has been tested using a single typical noise path between 0 and 1, so the two answers would be expected to differ by 0.5. The two values obtained were -0.4347 (with an absolute error bound of 0.044) using the Itô calculus and a value of 0.008118 (with an absolute error bound of 0.012) for the Stratonovich calculus. This is a difference of 0.4428 with an absolute error bound of 0.056.

There is a transformation between the two different formulations (provided that some technical assumptions are fulfilled), and so the two forms are not unrelated and the co-existence of two different forms of calculus is not a problem.

The transformation between the two forms is brought about by applying a drift correction to the coefficients in equation (2). Defining

$$\tilde{a}(t, x) = a(t, x) - \frac{1}{2}b(t, x)\frac{\partial b}{\partial x}(t, x), \quad (9)$$

and using this term instead of a in the Stratonovich calculus gives the same problem as (2) in the Itô calculus. Then

$$(\text{Itô}) \quad dX(t) = a(t, X(t))dt + b(t, X(t))dW(t), \quad X(0) = X_0, \text{ and}$$

$$(\text{Stratonovich}) \quad dX(t) = \tilde{a}(t, X(t))dt + b(t, X(t)) \circ dW(t), \quad X(0) = X_0$$

will be equivalent. This correction can be obtained by applying Itô's formula (see section 10.1.3.1) to obtain stochastic Taylor expansions for the a and b terms in the Stratonovich calculus and then considering how terms tend to zero in the mean square sense.

As an example, the SDEs

$$(\text{Itô}) \quad dX(t) = \lambda X(t)dt + \mu X(t)dW(t), \quad X(0) = X_0 \quad (10)$$

$$(\text{Stratonovich}) \quad dX(t) = \lambda X(t)dt + \mu X(t) \circ dW(t), \quad X(0) = X_0 \quad (11)$$

do not have the same solutions. The drift corrected forms are

$$(\text{Itô}) \quad dX(t) = \lambda X(t)dt + \mu X(t)dW(t), \quad X(0) = X_0 \quad (12)$$

$$(\text{Stratonovich}) \quad dX(t) = \left(\lambda - \mu^2/2\right)X(t)dt + \mu X(t) \circ dW(t), \quad X(0) = X_0. \quad (13)$$

Both forms will have solution

$$X(t) = X_0 \exp\left\{\left(\lambda - \mu^2/2\right)t + \mu W(t)\right\}. \quad (14)$$

10.1.3 Useful formulae and extensions

10.1.3.1 Itô's Formula

Section 10.1.2 mentioned that the chain rule for deterministic functions applies in Stratonovich calculus. It does not apply in Itô calculus, but an equivalent form called Itô's formula exists. Suppose that $Y(t) = f(t, X(t))$ is some function of the solution to an

equation of the form (2) and an equation for dY is required. Then Itô's formula gives

$$dY(t) = L^0 f(t, X(t))dt + L^1 f(t, X(t))dW(t)$$

$$L^0 = \left\{ \frac{\partial}{\partial t} + a \frac{\partial}{\partial x} + \frac{b^2}{2} \frac{\partial^2}{\partial x^2} \right\}, \quad L^1 = \left\{ b \frac{\partial}{\partial x} \right\}.$$

So, for instance, $f = X$ gets the form (2) back. An alternative way of writing this is

$$Y(t) = Y(t_0) + \int_{t_0}^t \left\{ \frac{\partial}{\partial s} + a \frac{\partial}{\partial X} + \frac{b^2}{2} \frac{\partial^2}{\partial X^2} \right\} f(s, X(s))ds + \int_{t_0}^t \left\{ b \frac{\partial}{\partial X} \right\} f(s, X(s))dW(s). \quad (15)$$

Repeated applications of the formula give stochastic Taylor expansions. For instance, consider the Stratonovich calculus:

$$dX_i = a(t_i + \Delta t/2, X(t_i + \Delta t/2))\Delta t + b(t_i + \Delta t/2, X(t_i + \Delta t/2))dW_i. \quad (16)$$

Applying Itô's formula gives

$$da = \left\{ \frac{\partial}{\partial t} + a \frac{\partial}{\partial x} + \frac{b^2}{2} \frac{\partial^2}{\partial x^2} \right\} a(t_i, X(t_i)) \frac{\Delta t}{2} + \left\{ b \frac{\partial}{\partial x} \right\} a(t_i, X(t_i)) \frac{dW_i}{2} \quad (17)$$

$$db = \left\{ \frac{\partial}{\partial t} + a \frac{\partial}{\partial x} + \frac{b^2}{2} \frac{\partial^2}{\partial x^2} \right\} b(t_i, X(t_i)) \frac{\Delta t}{2} + \left\{ b \frac{\partial}{\partial x} \right\} b(t_i, X(t_i)) \frac{dW_i}{2} \quad (18)$$

Inserting these expressions into (15) gives an expression in varying powers of Δt and dW_i . In the mean square sense, terms in $(dW_i)^2$ decay to zero like Δt , and so such terms join the drift coefficient so that (16) can be written

$$dX_i = \left(a(t_i, X(t_i)) + \frac{b}{2} \frac{\partial b}{\partial x} \bigg|_{t_i, X(t_i)} \right) \Delta t + b(t_i, X(t_i))dW_i$$

and so it is clear that the drift correction (9) will make this expression the same as that for the Itô calculus for a first approximation. Re-insertion of the integral forms of (17) and (18) into (3) gives

$$X(t) = X(t_0) + a(t_0, X(t_0)) \int_{t_0}^t ds + b(t_0, X(t_0)) \int_{t_0}^t dW(s) + \int_{t_0}^t \int_{t_0}^s L^0 a(\tau, X(\tau)) d\tau ds$$

$$+ \int_{t_0}^t \int_{t_0}^s L^1 a(\tau, X(\tau)) dW(\tau) ds + \int_{t_0}^t \int_{t_0}^s L^0 b(\tau, X(\tau)) d\tau dW(s) + \int_{t_0}^t \int_{t_0}^s L^1 b(\tau, X(\tau)) dW(\tau) dW(s) \quad (19)$$

as the stochastic Taylor expansion for $X(t)$. Higher-order terms can be obtained by repeating this process.

10.1.3.2 Fokker-Planck equations

The Fokker-Planck equations are used to determine the transitional probability distribution for a process of the form (2). The transitional probability distribution answers the question "Given the current position, what is the probability of being in some other region at some other time?": $G(y, t | x, s) = P(X(t) \leq y | X(s) \leq x)$. This would be useful for safety testing: if there is some critical value y that a quantity X must not

exceed, and it is known that the critical value is not exceeded at some time s , the transitional probability can give the probability that the critical value is not exceeded at a later time t .

It can be shown that if the drift a and the diffusion b are sufficiently smooth, the transitional density function g (the derivative of G with respect to y) satisfies

$$\frac{\partial g}{\partial t} + \frac{\partial}{\partial y}(a(t, y)g) - \frac{1}{2} \frac{\partial^2}{\partial y^2}(b^2(t, y)g) = 0, \quad (s, x) \text{ fixed}$$

which is called the Fokker-Planck equation or the forward Kolmogorov pde. A similar equation exists for the reverse case with (t, y) fixed, called the backwards Kolmogorov pde:

$$\frac{\partial g}{\partial s} + a(s, x) \frac{\partial g}{\partial x} + \frac{1}{2} b^2(s, x) \frac{\partial^2 g}{\partial x^2} = 0, \quad (t, y) \text{ fixed}.$$

For a standard Wiener process, $a = 0$ and $b = 1$, so the equations are a forward and a backward heat equation.

10.1.3.3 Multi-dimensional problems

The foregoing arguments have all been based on a single quantity, but the methods can be applied to multi-variate problems. A multi-dimensional Wiener process is written $\mathbf{W}(t) = \{W_1(t), W_2(t), \dots, W_m(t)\}$, and the components of this vector are pair-wise independent Wiener processes. Then the N -dimensional Itô diffusion process $\mathbf{X}(t) = \{X_1(t), X_2(t), \dots, X_N(t)\}$ satisfies

$$d\mathbf{X}(t) = \mathbf{a}(t, \mathbf{X}(t))dt + \sum_{j=1}^m \mathbf{b}_j(t, \mathbf{X}(t))dW_j(t)$$

where $\mathbf{a}$ and the $\mathbf{b}_j$ are N -dimensional vector functions. The components of the drift correction (5) become

$$\tilde{a}_i(t, \mathbf{x}) = a_i(t, \mathbf{x}) - \frac{1}{2} \sum_{j=1}^m \sum_{k=1}^N b_{k,j}(t, \mathbf{x}) \frac{\partial b_{i,j}}{\partial x_k}(t, \mathbf{x}), \quad i = 1, 2, \dots, N,$$

and it is clear that the formulae are the same for $m = N = 1$. Itô's formula also changes in more than one dimension. If $Y(t) = f(t, \mathbf{X}(t))$ is a scalar function, the rule becomes

$$dY(t) = L^0 f(t, \mathbf{X}(t))dt + \sum_{j=1}^m L^j f(t, \mathbf{X}(t))dW_j(t)$$

where

$$L^0 = \frac{\partial}{\partial t} + \sum_{k=1}^N a_k \frac{\partial}{\partial x_k} + \sum_{k,l=1}^N \sum_{j=1}^m \frac{b_{k,j} b_{l,j}}{2} \frac{\partial^2}{\partial x_k \partial x_l},$$

$$L^j = \sum_{k=1}^N b_{k,j} \frac{\partial}{\partial x_k}, \quad j = 1, 2, \dots, m$$

This further complicates the stochastic Taylor expansions, which will not be detailed here. In this notation, the noise term is called **commutative** if $L^j b_{i,j} = L^K b_{i,K}$ for all

$i = 1, 2, \dots, N$, and all $J, K = 1, 2, \dots, m$. These relationships can be used to simplify numerical schemes as they remove the need to calculate mixed integrals.

10.1.4 Numerics for SDEs

When solving SDEs numerically, it is tempting to simply use a scheme developed for deterministic equations. However, great care must be taken - often such a scheme may not be convergent, as some common schemes (e.g. Heun's algorithm, a simple predictor-corrector) are not consistent with the stochastic calculus. Additionally, high order deterministic schemes may not be of high order when applied with stochastic forcing and so will not offer any computational advantage over simpler schemes. This necessitates development of new numerical schemes.

The simplest scheme is the Euler-Maruyama method, which is derived by truncating the probabilistic Taylor series expansion (19) after the first three terms. Consider the process

$$X(t) = X_0 + \int_0^t a(s, X(s))ds + \int_0^t b(s, X(s))dW(s), \quad 0 \leq t \leq T, \quad (20)$$

as given in (3). This can be approximated by dividing the time interval into N equally sized finite steps, with $N\Delta t = T$, so that (20) over each step becomes

$$X(t_i) = X(t_{i-1}) + \int_{t_{i-1}}^{t_i} a(s, X(s))ds + \int_{t_{i-1}}^{t_i} b(s, X(s))dW(s).$$

This form is then approximated by writing

$$X(t_i) = X(t_{i-1}) + a(t_{i-1}, X(t_{i-1}))\Delta t + b(t_{i-1}, X(t_{i-1}))(W_i - W_{i-1}), \quad i = 1, 2, \dots, N,$$

where the noise increments $W_j - W_{j-1}$ have mean zero and variance Δt . This form is consistent with the Itô calculus because the b term is evaluated at the left-hand end of the interval. This is the Euler-Maruyama scheme, also called the stochastic Euler scheme. It is equivalent to the deterministic Euler method as mentioned in the report on Model Validation in Continuous Modelling [1]. Like the deterministic scheme, the stochastic scheme has poor stability properties.

Other methods exist for solution of these integrals. The methods, and issues connected with them, will only be summarised here, but further details can be found in Kloeden and Platen [35]. In general, higher order schemes can be found from the stochastic Taylor expansion (19) by including more terms or generating further ones using the Itô formula. For example, the Milstein scheme,

$$X(t_i) = X(t_{i-1}) + a(t_{i-1}, X(t_{i-1}))\Delta t + b(t_{i-1}, X(t_{i-1}))(W_i - W_{i-1}) + \frac{1}{2} b(X_{j-1}) \frac{\partial b}{\partial x} [(W_j - W_{j-1})^2 - \Delta t],$$

includes the final term in equation (19) which turns out to be the next lowest order term in Δt . The inclusion of this term gives the scheme strong order of convergence 1 (see section 10.1.5 for an explanation of convergence). Note that for additive noise, b is constant and so this scheme is identical to the Euler-Maruyama scheme. Special properties of noise can often be used to simplify methods like this.

Higher order schemes are rarely used because their generation requires the approximation of multiple stochastic integrals, which is a difficult task. In general the Euler-Maruyama scheme gives reasonable results, in particular when the drift and

diffusion coefficients are nearly constant, but the scheme is only of low order and higher order schemes are more accurate.

Implicit schemes are sometimes used, but they need careful derivation in order to be consistent with the Itô calculus. Often the best way is to use an explicit formulation for the stochastic noise term and an implicit formulation for the deterministic drift term so that any problems with stiffness in the drift term can be avoided. For instance, the semi-implicit stochastic Euler method can be written

$$X(t_i) = X(t_{i-1}) + a(t_i, X(t_i))\Delta t + b(t_{i-1}, X(t_{i-1}))(W_i - W_{i-1}),$$

with an implicit deterministic term and an explicit stochastic one. Derivative free schemes (also mentioned in Kloeden and Platen [35]) use intermediate steps to avoid approximating the derivatives directly, in a similar way to the multistep schemes of deterministic equations. These methods can have high orders of strong convergence without the need for calculation of difficult integrals.

Several methods have been designed to reduce the number of trials required to obtain an accurate estimate of quantities. Variance reduction techniques attempt to reduce the number of trials required by estimating a different quantity that has the same mean but a smaller variance. This idea is often used for Monte Carlo simulations.

Recent research has investigated adaptive step methods similar to those used to solve deterministic equations. To adapt the step it is often necessary to use a Brownian Bridge to get intermediate values of the noise. The method is explained more fully in Lamba et al [44].

Another issue for numerical schemes, particularly those involving the testing of strong convergence, is numerical noise generation. This requires a random number generator in order to sample from the appropriate distribution (a normal distribution for a Wiener process, as explained in section 10.1.1). Kloeden and Platen [35] mention various random number generators in chapter 1.3, and a companion document to the BPG on uncertainties and statistical modelling reviews the topic extensively for a variety of input distributions [45]. *Numerical Recipes* [27] also covers the topic.

10.1.5 Convergence and stability

When choosing a numerical scheme, convergence needs to be considered. The convergence compares a true solution $X(t_n)$ and an approximation X_n , both of which are random variables. It is usual to classify the convergence as **weak** or **strong**, depending on whether the information of interest is the statistics of the process (so that the probability distribution needs to be a good approximation to the real one) or whether the individual realizations of the process are required to be good approximations to reality. In metrological applications it is most likely that the first type of convergence will be of most interest.

Consider a fixed interval $[0, T]$ and let Δt be the fixed time step. Then a numerical scheme is of

- **strong order** γ if, for sufficiently small Δt , $E(|X(T) - X_N|) \leq K_T \Delta t^\gamma$. $E(|X(T) - X_N|)$ is the expected (mean) value of the error. Strong convergence asks whether individual resolution paths are approximated well, and looks at the mean of the error as $\Delta t \rightarrow 0$.
- **weak order** β if, for some function g and Δt sufficiently small, $|E(g(X(T))) - E(g(X_N))| \leq K_{\{g, T\}} \Delta t^\beta$. Generally g is taken to be the identity

function and so the weak convergence looks at the difference between the means of the true solution and approximation but it is straightforward to test convergence of other statistics.

Note that these definitions differ from the deterministic definitions [1] due to the probabilistic nature of the process making the evaluation of a mean necessary.

Testing for either weak or strong convergence makes it necessary to estimate an expectation, and this is done by performing a number of runs and calculating an expectation from the results. Expectations are produced for several different time steps in order to estimate the order of convergence. Getting a good estimate of the expectation requires a large number of runs, and to test strong convergence each time step needs to use the same realization of the noise and hence sample path of the variable. This means that the smallest time step must be run first, and the solution can then be re-sampled to produce runs with larger time-steps for comparison. Figure 23 gives an example of this process, showing the re-sampling of a numerical solution to equation (12) to examine the behaviour with a time step 10 times larger. Weak convergence only concerns the means of the results from the computed paths and so it does not matter if different realizations of the noise are taken as Δt is changed.

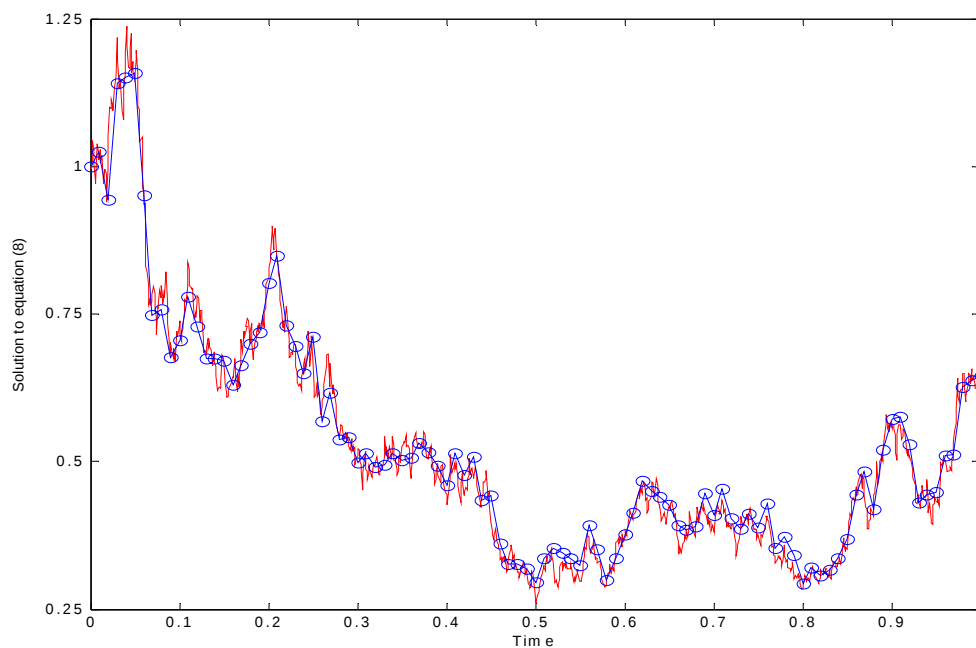


Figure 23: Numerical solutions to equation (12) calculated using the Euler-Maruyama method with time steps of 10^{-2} (blue circles) and 10^{-3} (red line). The solution with time step 10^{-2} is a re-sampling of the solution with time step 10^{-3} .

In general the stochastic Euler-Maruyama scheme has strong order $\gamma = \frac{1}{2}$. In the case of additive noise the strong order is improved to $\gamma = 1$. This is illustrated numerically in figure 24a. The stochastic Euler-Maruyama scheme has weak order $\beta = 1$, as shown in figure 24b. Different sample paths have been used for each Δt when demonstrating weak convergence, since averaged quantities are of interest and so any increment $W_j - W_{j-1}$ with appropriate statistical properties can be used.

When calculating weak convergence the increment can be replaced by an independent two-point random variable (as stated in section 10.1.1). Replacing the Wiener process in this way leads to the **Weak Euler-Maruyama** scheme which has weak convergence

of order $\beta = 1$. This scheme is not convergent pathwise and so has no strong convergence.

To prove strong convergence, it needs to be shown that a strong solution (i.e. one which is pathwise convergent to the solution of the SDE) for the SDE and the scheme exist. Under some assumptions on the measurability of $a(t, x)$ and $b(t, x)$, some Lipschitz assumptions on a and b , some bounds on linear growth, and assuming a sensible initial condition is used, then the strong solution can be shown to be unique, with $\sup_{0 \leq t \leq T} E(|X(t)|^2) < \infty$.

In general, if all that is of interest is the statistics of the process, it is best to use a weak scheme since weak schemes involve fewer terms of the expansion (19) than the strong schemes of corresponding order, and so weak schemes are more efficient.

Stability is another issue that needs to be addressed for numerical schemes, just as it does for deterministic methods. For stiff systems the standard Euler-Maruyama scheme is unstable and requires a restriction on the step size. Better stability can be obtained by using implicit methods, which are described in Kloeden and Platen [35] and Higham et al [46].

As an example of convergence testing, consider applying the Euler-Maruyama scheme to the linear SDEs

$$dX(t) = \lambda X(t)dt + \mu X(t)dW(t), \quad X(0) = X_0, \text{ and} \quad (21)$$

$$dX(t) = \lambda X(t)dt + \mu dW(t), \quad X(0) = X_0 \quad (22)$$

with $\lambda = -0.5$, $\mu = 0.25$, and $X_0 = 1.0$.

Figures 24a and 24b show the weak and strong convergence properties of the Euler-Maruyama scheme applied to equations (21) and (22). The figure shows logarithmic plots of $E(|X(T) - X_N|)$ in (24a) and $|E(X(T)) - E(X_N)|$ in (24b), plotted against time step in both cases. The lines without symbols are guide lines of slopes 1.0 (red) and 0.5 (green) for comparison purposes.

The slopes of the best fit straight lines to these data give the convergence orders for the two types of convergence described above. The strong convergence properties for the two equations differ. The best fit straight line to the approximation to equation (21), the equation with multiplicative noise, has a slope of 0.51 whereas the fit to the approximation to equation (22) and its additive noise has a slope of 0.90. These compare well to the expected values of 0.5 and 1.0 respectively. The weak convergence properties are similar for the two equations, with best fit straight line slopes of 1.00 for the additive noise and 0.98 for the multiplicative noise, although the fit for the multiplicative noise is not that good.

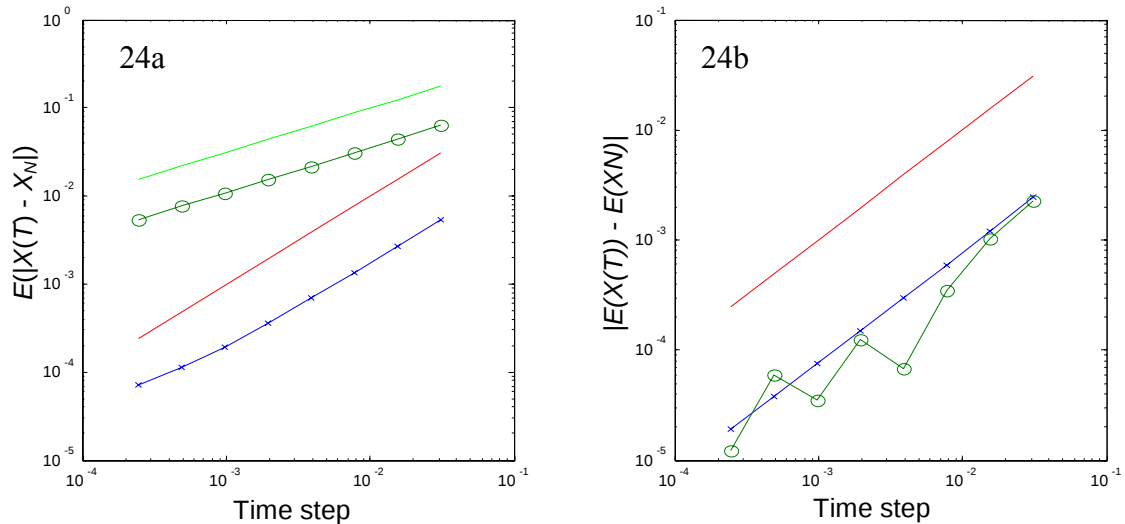


Figure 24: Plots showing the orders of strong (24a) and weak (24b) convergence for the Euler-Maruyama method applied to equations (21) (green line with circles) and (22) (blue line with crosses). Figures are log-log plots of $E(|X(T) - X_N|)$ (24a) and $|E(X(T)) - E(X_N)|$ (24b), plotted against time step. The lines without symbols are guide lines of slopes 1.0 (red) and 0.5 (green) for comparison purposes. The best straight line fit slopes are given in the text.

10.1.6 Further reading

The preceding sections have attempted to present the basic theory of stochastic differential equations, and to outline some of their properties and applications. Inevitably, much has been omitted or skimmed over, but there are a large number of books that go into detail and explain the material more thoroughly. Some further material that may be of interest includes:

- Kloeden and Platen [35]: an introduction to the general theory, and extensive demonstration of the methods required to solve SDEs numerically. These authors have also produced several other useful references [47, 48].
- Two commonly-used introductory texts by Gard [49] and Oksendal [50], the latter being particularly common for financial applications.
- Higham's article [51] and the accompanying Matlab scripts [52] have been adapted for use in this report.
- Cyganowski et al [53] have reviewed the use of MAPLE for the stochastic calculus.

This is by no means an exhaustive list, but is meant to indicate good introductory texts and software sources.