

NPL Report
DEM-ES 009

Software Support for Metrology
Good Practice Guide No. 13

Data Visualisation

John Gilby, Sira Ltd.
and Jeremy Walton, NAG Ltd.

NOT RESTRICTED

June 2006



The DTI drives our ambition of 'prosperity for all' by working to create the best environment for business success in the UK. We help people and companies become more productive by promoting enterprise, innovation and creativity.

We champion UK business at home and abroad. We invest heavily in world-class science and technology. We protect the rights of working people and consumers. And we stand up for fair and open markets in the UK, Europe and the world.

The National Physical Laboratory is operated on behalf of the DTI by NPL Management Limited, a wholly owned subsidiary of Serco Group plc

Software Support for Metrology

Good Practice Guide No. 13

Data Visualisation

John Gilby, Sira Ltd.
and Jeremy Walton, NAG Ltd.
with Trevor Esward, Peter Harris and Louise Wright
Mathematics and Scientific Computing Group

June 2006

ABSTRACT

This Guide gives information on good practice in the use of visualisation techniques applied to metrology data. It has been prepared to help metrologists and others who are concerned with the generation and interpretation of numerical data in the effective use of visualisation techniques. The aim of the Guide is to show the benefits of visualisation, give practical advice on how to perform visualisation, and to make existing good practice available in a readily accessible form.

© Crown copyright 2006
Reproduced with the permission of the Controller of HMSO
and Queen's Printer for Scotland

ISSN 1744-0475

National Physical Laboratory
Hampton 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.

We gratefully acknowledge the financial support of the UK Department of Trade and Industry (National Measurement System Directorate)

Approved on behalf of the Managing Director, NPL
by Jonathan Williams, Knowledge Leader for the Electrical and Software team

Contents

1	Introduction	1
1.1	Structure of the Guide	1
1.2	Background	2
1.3	What is visualisation?	3
1.4	The uses of visualisation	4
1.5	Why visualisation is now more accessible	5
1.6	Types of visualisation	5
1.7	Good practice in the visualisation of metrology data	6
1.8	Issues in visualisation arising from the increased use of computers	7
1.9	Recommendations	7
1.10	Acknowledgments	7
2	Requirements for Visualisation	9
2.1	The objective of the visualisation	9
2.2	Constraints on the delivery of the visualisation	10
2.3	Things to think about	10
2.4	Recommendations	12
3	Designing the Visualisation	13
3.1	The types of quantity being visualised	13
3.2	Visualisation dimensions	15
3.2.1	Spatial visualisation dimensions	15
3.2.2	Colour	28
3.2.3	Glyphs	32
3.2.4	Displaying edge- and face-based data	38
3.2.5	Animations and interactivity	39
3.2.6	Parallel co-ordinates	40
3.2.7	Combining visualisation dimensions	43
3.3	Mapping the quantity being visualised to a visualisation dimension	47
4	Implementation of the Visualisation	51
4.1	Development of visualisation software	51
4.2	Validation of visualisations	54
4.3	Recommendations	56
5	Visualisation of Uncertainties	57
5.1	Objective of the visualisation	57
5.2	Designing the visualisation	58

5.2.1	Co-ordinate system transformation	59
5.2.2	Slices	60
5.2.3	Elevated surfaces	61
5.2.4	Movies: slices, elevated surfaces and contour plots	63
5.2.5	Scatter plots	63
5.2.6	Scaled histograms	63
5.2.7	Interactive visualisation	65
5.2.8	Parallel co-ordinates	67
5.3	Recommendations	69
6	Visualisation in Continuous Modelling	71
6.1	Continuous modelling results	71
6.2	Interpolation, extrapolation, and smoothing	72
6.2.1	Interpolation and extrapolation	72
6.2.2	Smoothing	73
6.3	Visualisation parameter choices	74
6.3.1	Choice of viewing surface	76
6.3.2	Choice of local or global co-ordinate systems	76
6.3.3	Scaling factors	78
6.3.4	Display tolerances	78
6.4	Uncertainties and continuous modelling	79
6.4.1	Choice of output quantity	79
6.4.2	Choice of representation of the uncertainty	80
6.4.3	Software for visualising uncertainties associated with continuous models	80
6.5	Recommendations	81
7	Examples	82
7.1	Visualising models of experimental data	82
7.1.1	Introduction	82
7.1.2	Visualisations	83
7.1.3	Summary	91
7.2	Visualising uncertainties associated with the calibration of microwave power meters	91
7.2.1	Introduction	91
7.2.2	Visualisations	92
7.2.3	Summary	98
7.3	Visualising measurements of a sphere	99
7.3.1	Introduction	99
7.3.2	Visualisations	99
7.3.3	Summary	101
7.4	Visualising the results of continuous modelling	102
7.4.1	Introduction	102
7.4.2	Visualisations	102
7.4.3	Summary	107
7.5	Visualising ordinal and categorical data	107
7.5.1	Introduction	107
7.5.2	Visualisations	108
7.5.3	Summary	111
7.6	Validating a visualisation	113
7.6.1	Introduction	113

7.6.2	Visualisations	113
7.6.3	Summary	116
8	Summary	117
	Bibliography	120
A	Glossary	124
B	Multimedia objects	127
C	Further reading and useful resources	129
C.1	Textbooks	129
C.2	Web-based resources	131

Chapter 1

Introduction

This is one of a series of Best-Practice and Good-Practice Guides produced by the National Measurement System Software Support for Metrology (SSfM) programme.

This Guide has been prepared to help metrologists and others who are concerned with the generation and interpretation of numerical data in the effective use of visualisation techniques. The aim of the Guide is to show the benefits of visualisation, give practical advice on how to perform visualisation, and to make existing good practice available in a readily accessible form.

Visualisation in metrology is not new – almost all people involved in the handling and understanding of numerical data use graphical techniques. However, software tools currently available can give a much broader range of options for visualisation than were available in the past. In many cases those options are not explored because they are perceived to require too much effort for too little gain, or more generally because people are unsure how to go about the process of visualisation. It is hoped that by addressing both these issues this Guide will increase the appropriate use of data visualisation by the metrology community.

1.1 Structure of the Guide

Experience in the use of visualisation varies widely and it has been important to structure this Guide to reflect this variety. Some readers will already be constructing and using interactive three dimensional visualisations of data, while others will more typically restrict themselves to traditional graphs. The Guide has been structured so that it can be used both to develop a clear understanding of the visualisation process and to act as an aide-memoire of good practice.

The Guide is structured as follows:

- The remainder of this chapter gives the background to the Guide, explains visualisation, and sets it in context. A key message of this section is that visualisations need to be properly designed.
- Chapter 2 clarifies the process of analysing the requirements for visualisation. It is essential to make clear the requirements and objectives for visualisation, because if the objective of the visualisation is not clear then the resulting visualisation is likely to be unclear or inappropriate. The chapter includes a list of things to think about when developing a visualisation.
- Chapter 3 describes how to design visualisations. This topic forms the bulk of the Guide. The starting point is to understand the characteristics of the data or quantity required to be visualised. The chapter describes different types of *visualisation dimension*, and discusses how the data or quantity can be mapped to a visualisation dimension.
- Chapter 4 addresses issues arising when implementing the visualisation, including the validation of a visualisation.
- Chapter 5 discusses the visualisation of uncertainties, a topic of particular concern in metrology.
- Chapter 6 considers visualisation in the context of continuous modelling.
- Chapter 7 describes some more substantive case studies than the examples that are used for illustrative purposes throughout the text.
- Chapter 8 contains a summary.
- Appendix A provides a glossary of terms used throughout the text.
- Appendix B describes the multimedia objects contained within the document, and the viewers required for the objects to be visible.
- Appendix C provides a list of useful resources for visualisation.

Throughout the Guide pitfalls to avoid are highlighted and recommendations to the reader are made.

1.2 Background

The objectives of the SSfM programme are to identify good practice in the use of mathematics and software in metrology, to promote the awareness and application of such good practice across the whole breadth of metrology covered by the National Measurement System, and to advance the state of the art in appropriate areas of mathematics and software engineering.

Within the first SSfM programme (1998 – 2001), the need to increase awareness of visualisation amongst the metrology community was identified. It was recognised that

visualisation was becoming increasingly important due to increases in computing power, the greater use of data modelling, and the availability of high-quality visualisation tools. However, it was also identified that the use of visualisation by the metrology community was patchy. Sira Ltd. undertook work to promote awareness of visualisation techniques, and to document a set of case studies drawn from different areas of metrology [22].

Within the second SSfM programme (2001 – 2004) Sira Ltd. and NAG Ltd. prepared the first edition of this Good Practice Guide on Data Visualisation [16], together with an associated training course based upon the Guide. Additionally, NPL published reports on the visualisation of uncertainty associated with classes of electrical measurements [9], and uncertainties and visualisation in continuous modelling [41].

This revised edition of the Guide has been produced as part of the third SSfM programme (2004 – 2007). The principal modifications to the Guide are:

- A discussion in greater depth of the visualisation of uncertainties.
- The inclusion of a description of the parallel co-ordinates technique for visualising multivariate data¹.
- A consideration of visualisation in the context of continuous modelling.
- The updating and extending of the number and scope of examples included in the Guide.
- The inclusion of additional material covering topics such as the validation of visualisations and useful resources for visualisation.

1.3 What is visualisation?

The ability to record data from the world or to generate numerical results from models is now greater than ever. The data can be extensive in its volume and in its complexity. In order to make use of that data it needs to be interpreted; in other words, it needs to be transformed into information.

Complex data can be explored in a form close to its raw state and the information gained may be presented in words and numbers. However, such an approach can fail to extract the full information because of difficulties discerning the important features within the data. This is where data visualisation can help. Visualisation is communicating complex information graphically. It can be used both for exploring data and for presenting the derived information to others in an accessible form.

Visualisation, in general, includes all forms of graphical communication of complex information. It is not limited to numerical data and it includes graphical forms that are very familiar to metrologists such as simple two-dimensional graphs.

It is important to note that transforming visualised data into information requires interpretation. This means visualisation is not mechanistic – it is a design process. In forming a visualisation the author is making decisions about what is important within the

¹This material is not included in this draft of the Guide but will be part of the final release.

data and how best to convey the information derived. Taking those decisions explicitly, rather than blindly following typical practice or relying on the defaults of a software package, can lead to a much greater understanding of the data.

While visualisation needs to be strongly grounded within the data, its success relies on the human visual system and human perception. As such, it is important that proper account is taken of the range of human perception capabilities and some care should be taken not just to design visualisations for their author or for the ‘average’ person.

1.4 The uses of visualisation

There are three main uses for visualisation: exploring data, presenting information, and creating attractive pictures. This Guide is concerned with the first two of these uses.

In exploring data, visualisation is used to understand the characteristics of a data set. In general the person using visualisation will not have a full understanding of the data and he will be using the visualisation to probe aspects of the data. This kind of visualisation involves rapidly putting together views of the data so that the user can identify features of interest. Speed and simplicity are normally the priority although where data is very complex more sophisticated visualisation may be necessary.

Using visualisation to explore data can be subdivided into the following scenarios:

Directly using tools to explore the data. Examples of this situation occur when someone has recorded a new type of data set or derived it from a model and they use visualisation tools to understand the data. Typically, this involves the use of basic tools that can be assembled to suit the problem in hand.

Designing tools for others to explore the data. Where experiments or models are run repeatedly to produce data sets for others, those data sets will have considerable commonality and it may be appropriate to assemble higher-level tools that the end users can employ to explore the data.

When directly using tools to explore data, it can be appropriate to sacrifice detail for speed; for example, to produce a graph with un-labelled axes for a quick look at a certain aspect of the data. However, great care has to be taken that where such a quick visualisation is retained for further use, it is placed on a sound basis. When designing visualisation tools for others to explore data then it is essential that those tools are properly designed and follow good practice.

The general process in using visualisation to explore data is first to gain a qualitative overview which can be used to identify structures in the data such as patterns, trends, and relationships. That overview is then used to guide more focussed exploration leading to quantitative analysis. Exploratory data visualisation can be particularly useful for identifying anomalies in data whether they are the result of mistakes or are unexpected features of underlying processes forming the data.

Once a clear understanding of the data has been formed, then visualisation can be used to present that understanding to others. In contrast to exploration, when presenting data it

is known which characteristics of the data are important to communicate. Similar tools and techniques are employed, but the overall objective is different.

Effective visual communication of complex information requires clarity, precision and efficiency. Clarity and efficiency are closely linked – if the reader is presented with the appropriate information in a clear and unambiguous form then he will quickly be able to understand and use the information presented. Precision in the metrology context is particularly important. The visualisation must not give a misleading impression of the true data.

Visualisation tools can be used to present clear and concise information, but they can also be used to create attractive pictures whose primary objective is not the presentation of information. Such pictures can make interesting and colourful additions to reports and other documents. However, while clarity, precision and efficiency generally do lead to attractive pictures, care has to be taken not to confuse attractiveness for content.

1.5 Why visualisation is now more accessible

The very rapid advances in information technology over the last three decades have made powerful computers with extensive graphics capabilities available to all at modest cost. Until 1995, high-quality graphical display of data required specialist Unix workstations – typically from manufacturers like Silicon Graphics (SGI). In the early 1990s, SGI were leaders in the promulgation of 3D graphics on the desktop, thanks both to their specialised graphics hardware (the so-called geometry engine) and to their software. The latter was originally based on a proprietary graphics library (called, simply, GL) that evolved into a new library called OpenGL which SGI licensed to other vendors. In 1995 Microsoft launched the 32-bit operating system Windows 95 and included OpenGL as part of it. The enormous success of Windows 95, including its 3D graphics platform, had an impact on peripheral hardware development (specifically for the PC games market), resulting in the wide range of powerful, cheap PC graphics accelerators which support OpenGL.

The hardware support for graphics is now ubiquitous. Today the standard computers that metrologists use are quite sufficient for fast visualisation of data. In addition, advances in development packages and desktop visualisation applications now make the creation of programs with 2D and 3D graphical output considerably simpler.

Those who do not currently exploit visualisation may be concerned that graphical presentation can produce a distorted appearance of the data; however, the growth in hardware and software support for graphics has been matched by the availability of high-quality visualisation tools. These make it possible to create concise, accurate graphical presentations of data.

1.6 Types of visualisation

Different types of visualisation can be distinguished:

Scientific visualisation is concerned with the exploration of data arising in science and

engineering. Usually the term is restricted to exploration rather than to the presentation of data. A classic example of scientific visualisation is the understanding of thunder cloud models.

Data visualisation is a more generic term than scientific visualisation. Data visualisation and scientific visualisation share the characteristic of focussing on numerical data but data visualisation is normally considered to encompass both exploration and presentation of information. In addition, data visualisation is also used in applications such as stock-market analysis, which are not normally considered part of science and engineering.

Information visualisation concerns the visualisation of non-numerical data. For example, visualisation of the relationships within a group of people falls within this category. Information visualisation is often used to understand the structure of abstract networks and the degree of closeness of objects or concepts.

Although different types of visualisation can be distinguished, there are many principles and techniques which they have in common. It is thus perhaps more useful to concentrate on the effective exploration of data and the presentation of information rather than on terminology.

Within data visualisation there are many application areas for which there is a distinct body of knowledge and expertise. Most notable is geographic data visualisation which has an extensive history due to the application of cartography. The other discipline which has used graphical presentation of information since long before computers became available is statistical analysis. Although both of these areas tend to use two-dimensional plots the expertise in them is widely applicable and they are primary sources of good practice in visualisation for metrologists.

1.7 Good practice in the visualisation of metrology data

Visualisation tools are now readily available to people who may not be familiar with how to use them effectively. While that ready availability is good, there is the risk of naive application leading to poor and misleading visualisation. More positively, good visualisation generally depends on knowing the tools and techniques available and on common sense. In this respect it is no different from the use of two-dimensional graphs.

For metrologists, it is now straightforward to explore and communicate data graphically; however, the opportunities to distort data and present it in misleading ways are now greater than ever. By employing existing good practice, we can make best use of our data and avoid misleading interpretations of it.

1.8 Issues in visualisation arising from the increased use of computers

Although the use of computers in visualisation has lead to many benefits, there are drawbacks.

- Many of the software packages are easiest to use by those who have a good computer science or programming background. In some packages this problem is reduced by having graphical programming interfaces, but in general this can be an obstacle for some users.
- Within any visualisation toolkit, the range of tools and their ease of use influence the kind of exploration and presentation that it is natural to perform. This is particularly problematic with packages such as Microsoft Excel which are not primarily designed for data visualisation.
- The tool in which a visualisation is produced is often also the one with which it is most convenient to view and manipulate the visualisation. This can make it difficult or inconvenient to share visualisation with others who do not have similar tools.

Sensible choice and use of visualisation tools can avoid these drawbacks.

1.9 Recommendations

- The visualisation of metrology data is a design process.
- Have a clear view of what you are trying to achieve when you use data visualisation. Are you exploring data, communicating it to others, or just creating pretty pictures?
- Use appropriate tools.
- Select visualisation techniques on the basis of their effectiveness as well as their availability.
- Check that your visualisations are clear, precise, and efficient.

1.10 Acknowledgments

This guide was produced under the SSfM programme, which is managed on behalf of the DTI by the National Physical Laboratory. For more information on the SSfM programme visit the SSfM website at www.npl.co.uk/ssfm or contact the programme manager Mr Bernard Chorley on +44 20 8943 7050 (email: ssfm@npl.co.uk) or by post to National Physical Laboratory, Hampton Road, Teddington, Middlesex, UK TW11 0LW.

The figures in this Guide have predominantly been prepared using the software packages IRIS Explorer and MATLAB. Further information on these software tools may be obtained from:

IRIS Explorer

NAG Ltd
Wilkinson House
Jordan Hill Road
Oxford
UK, OX2 8DR
Telephone: +44 1865 511245
Fax: +44 1865 310139
E-mail: sales@nag.co.uk
Web: www.nag.co.uk

MATLAB

The MathWorks Ltd
Matrix House
Cowley Park
Cambridge
UK, CB4 0HH
Telephone: +44 1223 423 200
Fax: +44 1223 423 289
E-mail: info@mathworks.co.uk
Web: www.mathworks.co.uk

Most of the parallel co-ordinate plots in sections 3.2.6 and 5.2.8 were created using Curvaceous Software's CVE. See www.curvaceous.com for more detail.

The data shown in Figure 3.10 was supplied by partners within the REVEAL project. REVEAL is a research and technological development project funded by the European Commission under the Competitive and Sustainable Growth Programme (Project no: 1999-RD-10657).

Figures 3.19, 3.20(a) and 3.25 are based on data supplied by Dr John Redgrove, NPL.

Example 1 (section 7.1) uses data supplied by Dr Mike Downs, NPL.

Section 5 and Example 2 (section 7.2) is based on joint work with NPL colleagues Prof Maurice Cox, Dr N McCormick, Mr Nick Ridler and Mr Martin Salter.

Figures 3.17 and 3.26 and Example 3 (section 7.3) use data supplied by Mr Alan Wilson, NPL.

Example 4 (section 7.4) uses data supplied by Mr Stephen Robinson, NPL.

Chapter 2

Requirements for Visualisation

For visualisation to be successful, it is important that the objectives of, and constraints on, the visualisation are carefully considered. Delivering a solution to the wrong problem or one which demands inappropriate time or resources to develop or use is worthless. Common mistakes made when using visualisation are not being clear why it is being used and not being able to judge whether it has been successful.

2.1 The objective of the visualisation

The first question to ask is ‘Why is visualisation being used for this data set?’ If the data set is small or otherwise simple then using graphical techniques may not be useful. If the data set is not simple then it is appropriate to consider who will be the end-users of the visualisation and what the visualisation will need to achieve for them if it is to be considered a success.

The approach to the visualisation may be different depending on whether it will be used for the one-off analysis of a data set, or used for a series of data sets, or be designed to give a general capability across a number of types of data. In general, the more extensive the flexibility desired, the greater will be not only the applicability but also the cost in design and execution time. It is recommended that a visualisation should be limited to particular data sets with high commonality. This approach makes it easier to create, and avoids the problem of speculating about data sets that never arise. Moreover, visualisation tends to be very suitable for rapid prototyping so it is sensible that complex visualisations can be developed iteratively from simple initial visualisations.

The other key question in formulating the visualisation requirements is ‘What is the important information that needs to be derived from the data?’ This helps to concentrate the process on the features and characteristics of the data that are useful to explore.

2.2 Constraints on the delivery of the visualisation

Where a visualisation is being built by the person who is using it for exploring a data set then he will inevitably have access to the tools to view and interact with the visualisation. This may not be the case when a visualisation is being built for others to use, and care has to be taken to ensure that the visualisation is accessible.

The simplest form of constraint is where the visualisation is dependent on a particular software application for viewing and manipulation. It can be a waste of time to develop an interactive visualisation if the users do not have the tools to interact with it. A typical solution to this problem is to deliver the visualisation in a much simpler form such as a static two-dimensional image. Where this is necessary, it is important that the simpler form conveys the necessary information, otherwise the purpose of the visualisation is lost.

When delivering the visualisation in a different form to the original, in general the interaction with the data will be reduced. For example, an interactive three-dimensional model may enable the user to zoom in onto particular features and refer back to the underlying data, whereas a static picture of the same model may only give qualitative information. In such situations, it is important to consider whether the end user will gain useful information from the simpler form.

A visualisation may also be dependent on the computer hardware available to the end user (memory, processor speed, graphics capability); although this constraint is now much less common (see section 1.5, above).

2.3 Things to think about

The following list of questions and responses may help to decide the requirements of, and constraints on, a visualisation.

- What is the purpose of the visualisation?
 - To explore experimental data.
 - To convey information to the reader of a paper or report.
 - To form part of a PowerPoint or other slide presentation.
 - To present calibration or measurement service results. For example, the user might want to read data from a graph to use in calibrating an instrument.
 - To compare two or more sets of data (e.g. theory with experiment).
 - To assist with some other activity, such as computational steering.
- What form will the visualisation take?
 - In a report or paper.
 - On slides.
 - On a computer screen.
 - Colour, grey scale or black and white.

- Static or dynamic presentation.
 - Interactive visualisation.
 - Stand-alone or with accompanying words.
- Who is the audience for the visualisation?
 - Yourself.
 - Peers in own branch of science.
 - Specialists
 - Mathematicians.
 - Theoreticians.
 - Experimentalists.
 - General scientific audience.
 - General non-expert audience.
- What is the nature of the data?
 - Categorical.
 - Ordinal.
 - Function of one, two, three, four or more variables.
 - Time history.
 - A geometrical object or a mathematical object.
 - Scalar-valued or vector-valued quantities.
 - Real-valued or complex-valued quantities.
 - At the micro- or nano-scale (e.g., molecular shapes and configurations)?
 - Uncertainties associated with underlying data?
 - Can the data dimensions be transformed to make them more meaningful, eg. by normalisation, finding dimensionless quantities, or exploiting symmetries?
- What visualisation software is available?
 - MATLAB.
 - MathCad.
 - Microsoft Excel.
 - IRIS Explorer.
 - Other (both freeware and proprietary).
- Do you need advice on appropriate visualisation software or links to downloadable software?
- Does the visualisation need to be validated?

Some of the possible answers to these questions, and the ways in which they affect visualisations, are discussed further in the rest of this Guide.

2.4 Recommendations

- Be clear about why visualisation is being used and what it has to achieve.
- Identify the important information which the visualisation will provide.
- Structure complex visualisations so that they can be built in stages.
- Provide the visualisation in the correct form for its intended end-use.
- Ensure that the correct balance of quantitative and qualitative information is provided.

Chapter 3

Designing the Visualisation

Information from visualisation is limited by human perception. We can identify characteristics of objects such as type, location, angle, size, colour hue, and colour saturation. Each of these characteristics can be used as a *visualisation dimension*. For example, colour hue may be used to represent temperature in a two-dimensional area plot of a temperature field. The basis for visualisation, therefore, is finding an effective mapping from the data being visualised to the visualisation dimensions. The difficulty in choosing the mapping arises for three principal reasons:

- The data may have many more dimensions than are available in the visualisation.
- The visualisation dimensions may be perceived in different ways.
- The visualisation dimensions may not capture all of the properties of the data dimensions.

Often it is necessary to break the visualisation into a series of views, each of which has its own mapping. This then creates the problem of how to relate information in the individual views.

The first step in designing a visualisation is to carefully consider the form of the data (section 3.1). This chapter also describes different types of visualisation dimension (section 3.2), including spatial, colour, glyphs and time, and how to combine visualisation dimensions. Also discussed is how to map the quantity being visualised to a visualisation dimension (section 3.3).

3.1 The types of quantity being visualised

The most basic classification of data types is into categorical, ordinal and continuous data.

Categorical data is discrete and unordered. Examples include which laboratory performed an experiment, to which experimental data set does a point belong, and to which batch does a tested part belong. The decision to treat the data as unordered is dependent on the

information that is important in a given context. For example, in a set of experiments we may want to identify only those that contain anomalous data. Alternatively, we may want to look at the trend in the experiments over time. Only in the second case are we considering the data to be ordered.

Ordinal data is discrete and ordered. An example is a set of experiments ordered by date. Data could also be *partially* ordered – for example, where a series of pair-wise comparisons are made but not all pairs are compared. Partially ordered data often occur in information visualisation and lends itself to being displayed using network diagrams.

Metrology data is usually *continuous*: its values are not selected from a discrete set and typically the values are expressed as floating point numbers. Although conceptually the distinction between ordinal and continuous data is precise, in practice, ordinal data may be treated as continuous. This may occur where we are counting events and the number becomes large or where we believe that a continuous variable is generating the discrete values.

The form of the data influences how it is appropriate to visualise it. Many visualisation techniques rely on implicit interpolation – categorical data cannot be interpolated and ordinal data can only be interpolated with care.

In addition to the classification into categorical, ordinal and continuous types, data may show more extensive structure. Besides *scalar*-valued quantities we may have *complex*-valued quantities (having real and imaginary parts) and *vector*-valued quantities (having multiple components). These quantities may have additional structure that we wish to preserve in the visualisation, for example, if we have a vector-valued quantity specifying position we may want a mapping to the visualisation dimensions that preserves the distance metric. Note that complex-valued quantities have a distance metric which can be usefully mapped onto spatial visualisation dimensions even if the quantity does not specify position in space.

Complex-valued quantities highlight one of the key problems in data visualisation. A single complex-valued quantity has a natural representation in two spatial dimensions. If we wish to discern the relationship between two complex-valued quantities then the natural representation would require four spatial dimensions whereas we are limited to directly experiencing only three. While it would be possible to use colour as a fourth visualisation dimension, we would lose the distance metric structure of one quantity and there is no natural perceptual mapping between distance and colour.

Vector-valued data may have structure beyond that of its individual components. The most common structure is when the data represents a directed quantity whose identity is independent of the co-ordinate system used. For example, the flow of heat at any point within a component has both magnitude and direction and heat flow vectors can be expressed with respect to an arbitrary co-ordinate system. This ability to express the data with respect to a change in the co-ordinate system means that the quantity is a tensor. Data may also occur as higher-order tensors, for example, engineering stress is a tensor quantity. The tensor quantities may also have complex-valued components, for example, polarised light can be described by two complex numbers in two orthogonal spatial directions which are perpendicular to the direction of travel.

It is important to preserve a sufficient amount of the internal structure of the data in the visualisation. This may be achieved by mapping the data onto visualisation dimensions

that have the same structure as the data or by using invariants of the data. For example, the first stress invariant (sum of principal stresses) is often used when visualising stress data because it is a scalar quantity which is physically meaningful.

In metrology the above discussion of data types applies to the values of the measured quantities. However, in addition, metrologists are fundamentally concerned with understanding and characterising uncertainties associated with measurements. This means that as well as the measurement, there is a need to visualise the associated uncertainty. Here, we are interested in probability distributions and the differences between measured (inexact) data and data produced by models. Again, this gives rise to many more dimensions within the data sets, and presents further challenges in the construction of the visualisation.

Recommendations

In considering the data to be visualised, it is important to consider the following.

- How many dimensions are there within the data?
- What structure is inherent in the data?
- What part of that structure should be preserved by the visualisation?

Generally, in non-trivial data visualisation applications there will be more data dimensions and structure than can be effectively displayed simultaneously.

3.2 Visualisation dimensions

3.2.1 Spatial visualisation dimensions

We live in a world with three spatial dimensions and the ability to perceive the positions of objects and their physical relationships is fundamental to human perception. This makes the spatial dimensions of the visualisation the primary way of conveying information.

Spatial dimensions are also important because data is often located spatially. For example, we may measure the stress at each position on the surface of an engineering component. Spatial dimensions admit a distance metric which may be used in the visualisation. In the stress example, we might directly map spatial dimensions in the world to those in the visualisation so that we are able to judge the distance between two data features. Almost all visualisation makes use of spatial visualisation dimensions, although there are exceptions such as where the data is mapped to another human sense such as hearing.

While visualisations of two-dimensional objects are common and easily understood, care needs to be taken with three- (or higher-) dimensional objects. Three-dimensional objects are almost always visualised in two-dimensions because computer screens and printed media are usually two-dimensional. This means that perceptual cues must be used to infer depth information. Figure 3.1(a) shows a data set plotted without such cues. It is not

possible to gain a useful sense of the data from this graph. Plotting the data as finite-sized points – as in Figure 3.1(b) – causes them to occlude each other (see below). Such a perceptual cue adds some sense of depth to the visualisation.

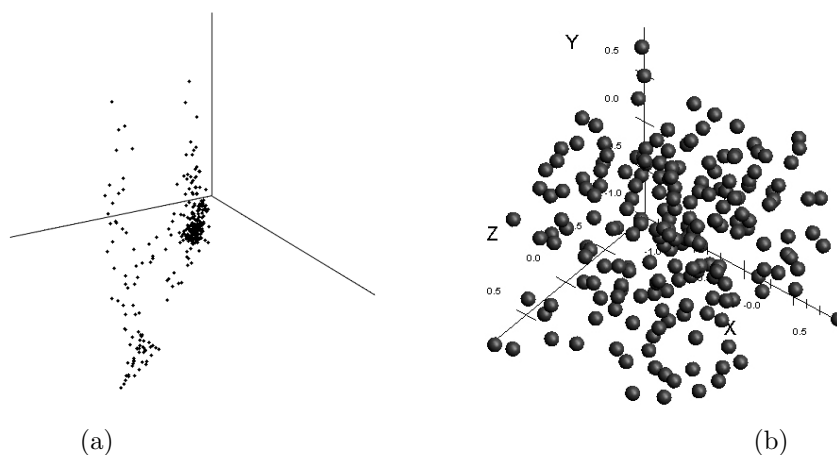


Figure 3.1: (a) Three-dimensional data flattened to a two-dimensional representation without the use of other perceptual cues. (b) A similar data set, but now plotted with finite sized points.

Alternatives to a two-dimensional display surface are stereoscopic displays and real three-dimensional objects. Stereoscopic displays are available, but they are not common and they only give limited perceptual cues to depth. Generally they give stereopsis – each eye seeing a slightly different image – but they do give vergence and focus cues. Although it is rarely justified to build an actual three-dimensional model, real objects give the best sense of three-dimensional structure. In almost all cases where such models exist, they are a model of a physical object such as a molecule rather than a model of experimental data. Note, however, that rapid prototyping techniques do exist for converting three-dimensional computer models into physical form.

Perceptual cues for inferring depth can be divided into those that are available in a static image and those which rely on the image changing. For a static image, the primary depth cues are occlusion, lighting and assumptions concerning orthogonality. When we look at Figure 3.2(a) we normally make the assumption that it is the projection of a rectangular parallelepiped. Figure 3.2(b) strengthens this perception by indicating occlusion information, and Figure 3.2(c) enforces the perception still further.

The assumptions we make concerning orthogonality are often described as perspective, whereas that term should more properly be used to refer to far objects appearing smaller than near objects. Perspective is very common in visualisation but should be avoided because it adds little to the perception of depth in diagrams and can lead to incorrect comparisons between different features in the visualisation – it can be difficult to distinguish object size variation from variation in closeness to the viewpoint. Figure 3.3 shows an identical graph with and without the use of perspective. The effect of perspective on the placement and size of the bars in the chart can be seen quite clearly. We note in passing that this use of a three-dimensional representation for a two-dimensional dataset –

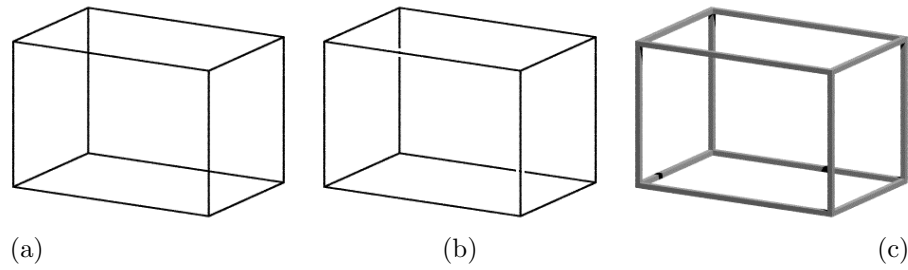


Figure 3.2: Increasing perceptual cues for inferring depth: (a) assumption of orthogonality, (b) addition of occlusion, and (c) use of lighting.

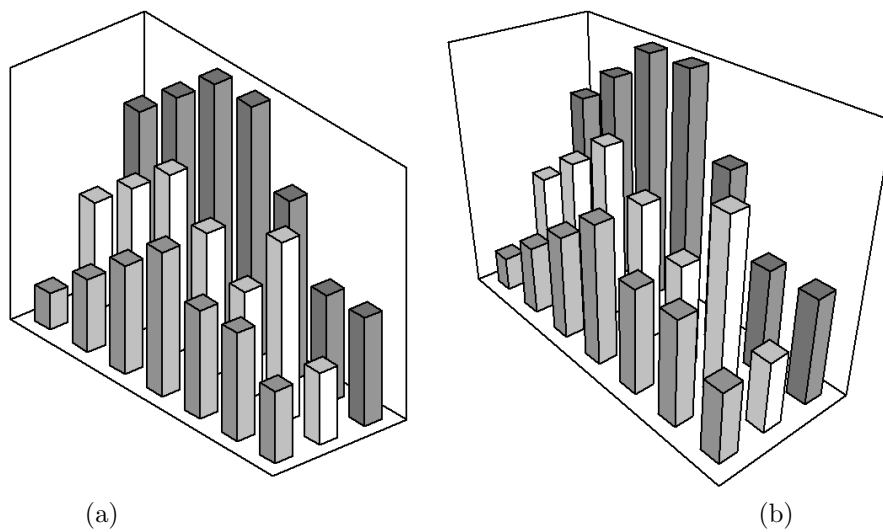


Figure 3.3: (a) No perspective (orthographic projection). (b) Perspective projection.

although widespread – could be regarded as gratuitous, since a two-dimensional representation would convey the same information without the ambiguity illustrated in Figure 3.3 – we return to this point below.

Occlusion can be both useful and problematic. Near objects occluding far objects gives a good sense of depth but it can also result in data being hidden. Careful choice of viewpoint can avoid the obscuration of important data. Figures 3.4(a) and 3.4(b) show the use of occlusion.

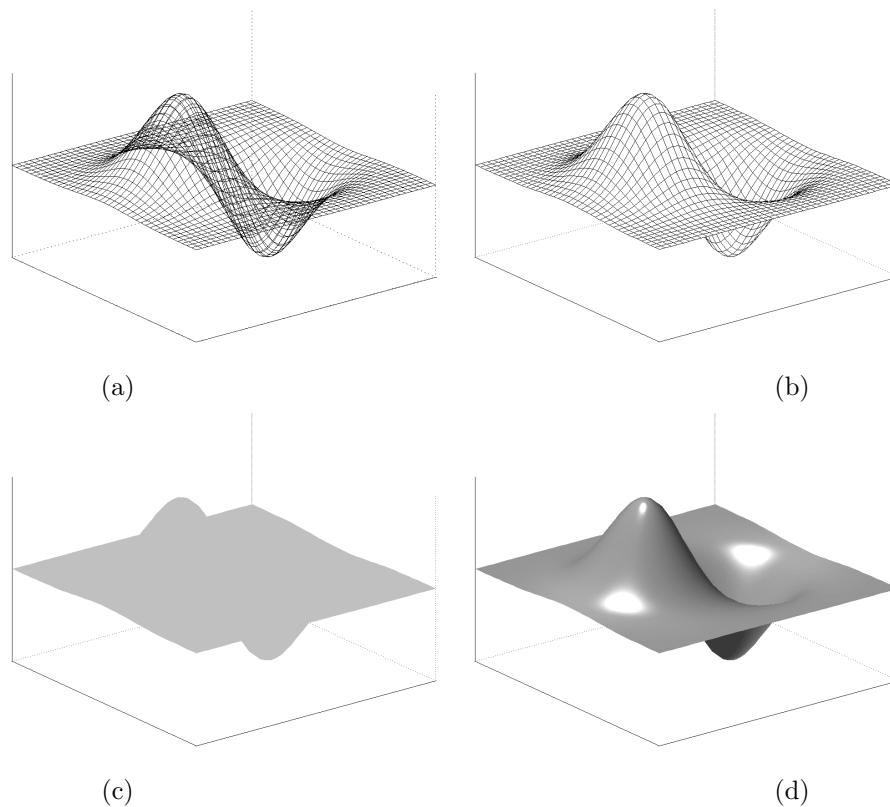


Figure 3.4: Use of occlusion and lighting depth cues.

Lighting can be a powerful cue to depth and it is widely used in computer interfaces and presentations to give a sense of three-dimensionality (shadowing on buttons, and drop shadows on text). Figures 3.4(c) and 3.4(d) show a simple function plotted with and without the use of lighting. The disadvantage of using lighting is that it may interfere with the use of colour as a visualisation dimension because it can alter saturation and hue.

The way in which the light interacts with a surface is dependent on the properties of the surface – whether it is rough, shiny, or glossy. These properties can be parameterised and their values changed by the user to affect the appearance of the surface. Figure 3.5 shows the same three-dimensional surface plotted with different lighting parameters. Figure 3.5(a) shows the surface with ambient and diffuse lighting, (b) and (c) show the same surface with increasing amounts of specular reflection – note the highlighted area caused by reflected light, and (d) shows the same surface as in (c) but now with decreased

shininess resulting in a highlight that is more spread out over the surface.

Figure 3.6 contains a similar sequence in which the light direction is manipulated to highlight different parts of the surface. The surface is illuminated by a light source at infinity (i.e., emitting parallel light rays) which is characterised by direction only. The figure shows various snapshots of the surface as the light direction is changed. Note that the lighting model only produces simple light and dark shading – there is no self-shadowing (caused by one part of the surface casting shadows on another, for example) in this model. Changing the light in this way, especially if it can be done interactively by the user, provides a rather direct way of exploring features in the visualisation, particularly properties such as surface curvature (see, e.g., Figure 3.6(c)).

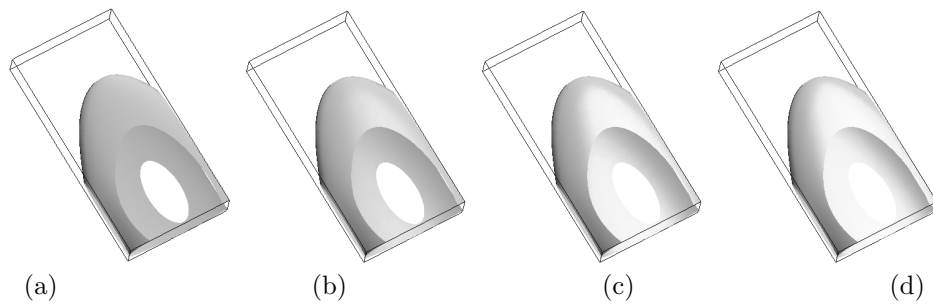


Figure 3.5: Effect of changing surface lighting parameters on visualisation appearance.

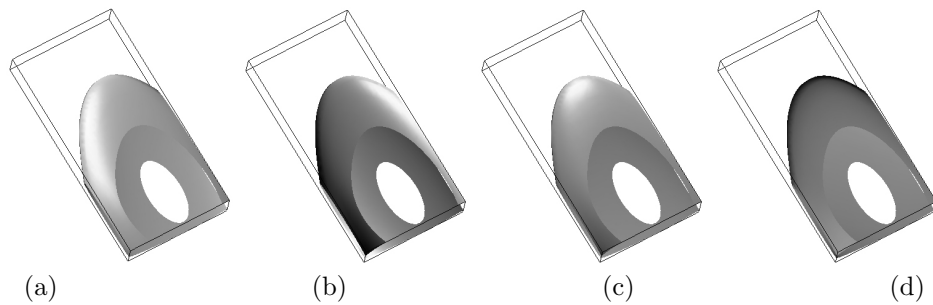


Figure 3.6: Effect of moving light source on visualisation appearance for the surface shown in Figure 3.5(d).

For non-static images, parallax is a powerful perceptual cue. If we are able to change the viewpoint of the projected image or it is animated then we immediately gain a stronger sense of its depth. Parallax cannot always be employed because the image may be static or movement may be used to indicate other information about the data. However, it should be used where possible to enhance the perception of depth in projections of three-dimensional models.

Being able to interactively change the viewpoint provides a deep insight into the structure or depth of a three-dimensional object. The usual model employed (see the sequence in Figure 3.7) is of a camera in three-dimensional space pointing at the object: the position of the camera can then be

- rotated (via a virtual trackball whilst retaining the point at which it is looking), as

shown in Figure 3.7(a), (b) and (c),

- panned (moved in the plane normal to its viewing direction), or
- dollied (moved closer to or further from the viewing point), as shown in Figure 3.7(d).

Some examples of interactive visualisations concerning the visualisation of uncertainty are provided in Chapter 5.

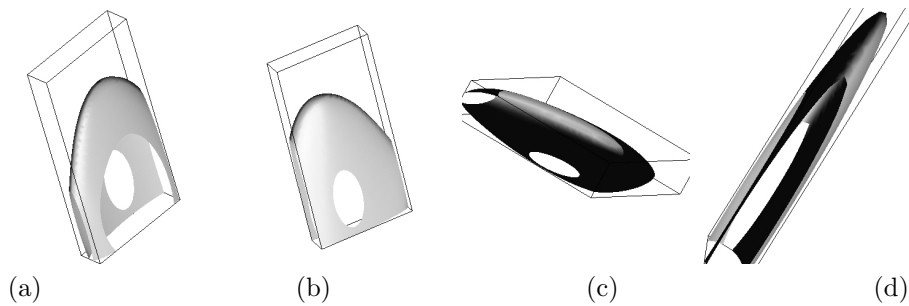


Figure 3.7: Effect of changing viewpoint on visualisation appearance for the surface shown in Figure 3.6(d).

Three-dimensional projections generally give a greater qualitative sense of the data at the expense of quantitative information. Figure 3.8 shows the same data plotted as a contour map (Figure 3.8(a)) and as a projection of a three-dimensional model (Figure 3.8(b)). It is much easier to gain quantitative information from the contour map, but most people find the projection easier to interpret rapidly. Note that perceptual cues such as lighting can be usefully employed on two-dimensional plots as shown in Figure 3.8(c), but care should be taken not to clutter plots with unnecessary perceptual cues.

The need to infer depth in projections of three-dimensional models can limit their usefulness and alternative ways of plotting the data should always be considered. However, where the data is inherently dependent on three spatial dimensions, the projection of three-dimensional models is invaluable. For example, a computational fluid dynamics model of flow past a three-dimensional object is difficult to visualise without the use of such projections.

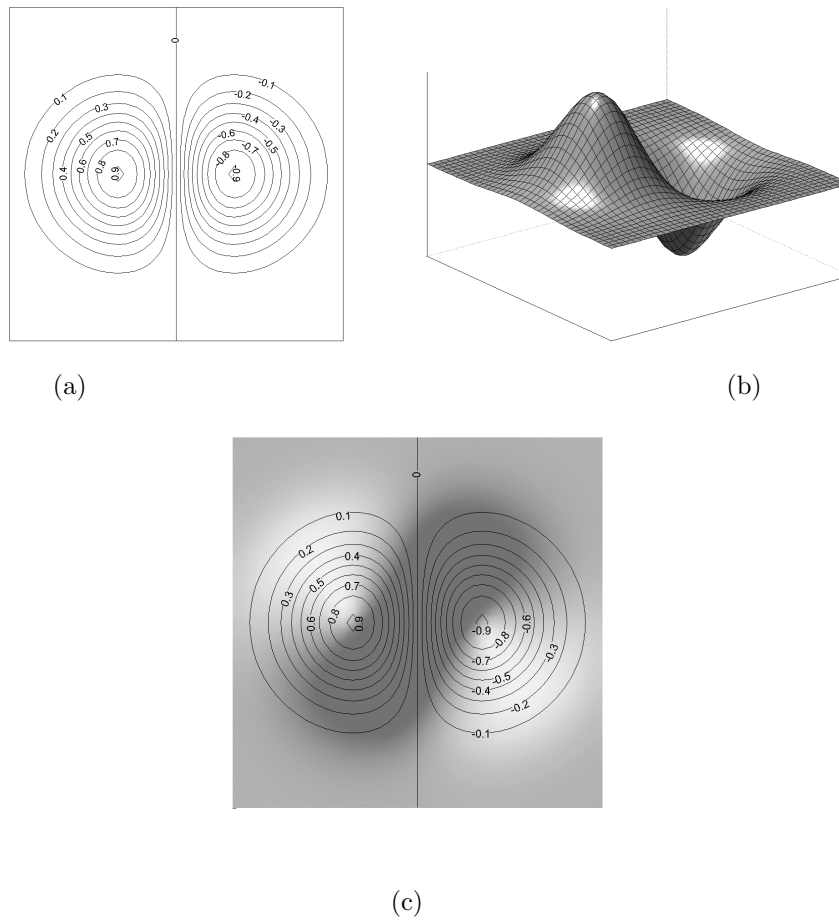


Figure 3.8: Two-dimensional contour maps, (a), can provide more quantitative information than three-dimensional projections, (b), but they can take more time to understand. Two-dimensional plots can be given lighting cues, (c).

The basic data that can be plotted spatially are points, lines, surfaces, and volumetric data. There are standard approaches to plotting each type of data. In two dimensions, point data has a unique position with respect to chosen co-ordinate axes but in projections of three dimensions further information is needed to assess depth. In a static image, this information can be provided by placing the point on a surface or by using drop lines (giving what are also known as stem plots). Drop lines are normally vertical as shown in Figure 3.9(a), but other directions and multiple drop lines can also be used as in Figure 3.9(b). Note that drop lines become confusing if the number of data points is too great (Figure 3.9(c)), unless there is clear structure within the data (Figure 3.9(d)).

Drop lines can also be used with lines, although in projections of three-dimensional objects it is more common to give them a physical form that can be lit – see section 3.2.2. Dashed lines should be avoided as they greatly increase the visual complexity of the data presentation, and shades of grey and colour should be used instead. Normally it is best to keep the number of lines on a plot to a minimum and to use multiple plots, particularly when quantities with different physical units are being displayed. As well as creating

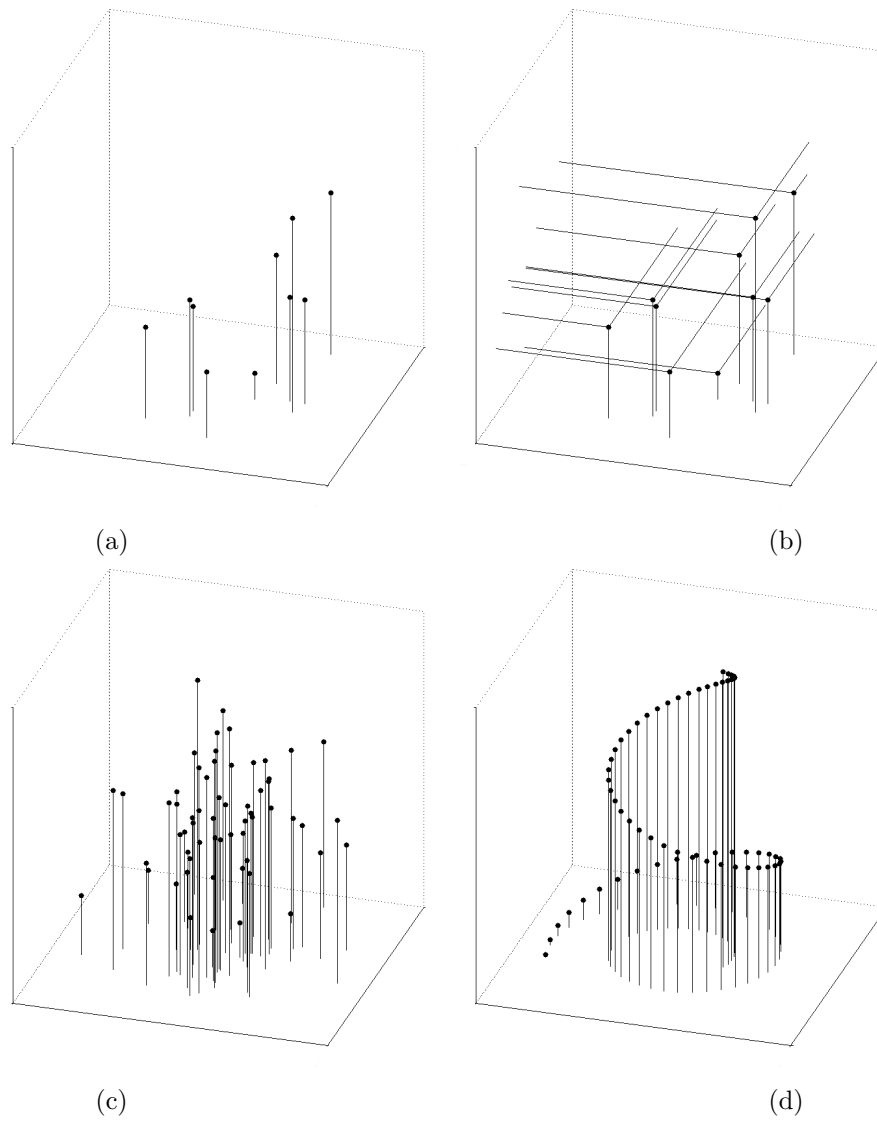


Figure 3.9: Drop lines can be used to indicate the positions of points in three-dimensional projections.

clearer diagrams, this avoids the use of legends which are commonly provided in visualisation packages but often do not add to overall legibility. Figure 3.10 shows vehicle exhaust emission data plotted on one graph and as multiple graphs.

In Figure 3.10 it is much easier to see in the lower graphs the correlation between the measurements of carbon dioxide, carbon monoxide and hydrocarbons. As these are vehicle tailpipe emissions, we would expect the ratio of these values to remain roughly constant over a short duration of a few tenths of a second. However, if we compare the individual plots of those channels we see the ratios in the first 0.06 seconds differ from the later values. This is not evident from the top graph, owing to the large variation in the range of the data.

Although point and line data can be plotted easily on paper and a computer screen, the display of surface data is more challenging, since it has as many dimensions as the display surface. Surface data can be defined either explicitly in the form $z = f(x, y)$ or implicitly in the form $f(x, y, z) = \text{constant}$. An explicit definition yields at most one z value corresponding to each (x, y) position, while an implicit definition can give multiple z values. The visualisation of each type of surface requires different techniques – thus, for example, Figure 3.8 shows examples of plotting an explicitly defined function as an elevated surface and a contour plot. Implicitly defined surface data cannot in general be plotted as a contour map. Generally, such data requires the use of three-dimensional projections (if the data cannot be separated into lower-dimensional projections) or animation. Figure 3.11 shows an implicitly defined surface plotted first with lighting cues and transparency and second with x and y contour lines. Note that the data is sufficiently simple for these plots to give a reasonable sense of the shape of the surface.

Volumetric data of the form $v = f(x, y, z)$ provides even more difficulties for plotting on a two-dimensional display surface. The principal techniques for visualising such data include reducing the spatial dimensions by taking two-dimensional slices through the data, plotting isosurfaces, and volumetric rendering.

An isosurface is the implicitly defined surface $v = f(x, y, z)$ with v held constant. Creating a series of isosurface plots for differing values of v (sometimes called the threshold value) can be used to build up a picture of the volumetric data set, as shown in Figure 3.12. In this example, the difference between successive threshold values is a constant, but other values could be chosen, particularly in regions of the parameter space where the function is changing rapidly.

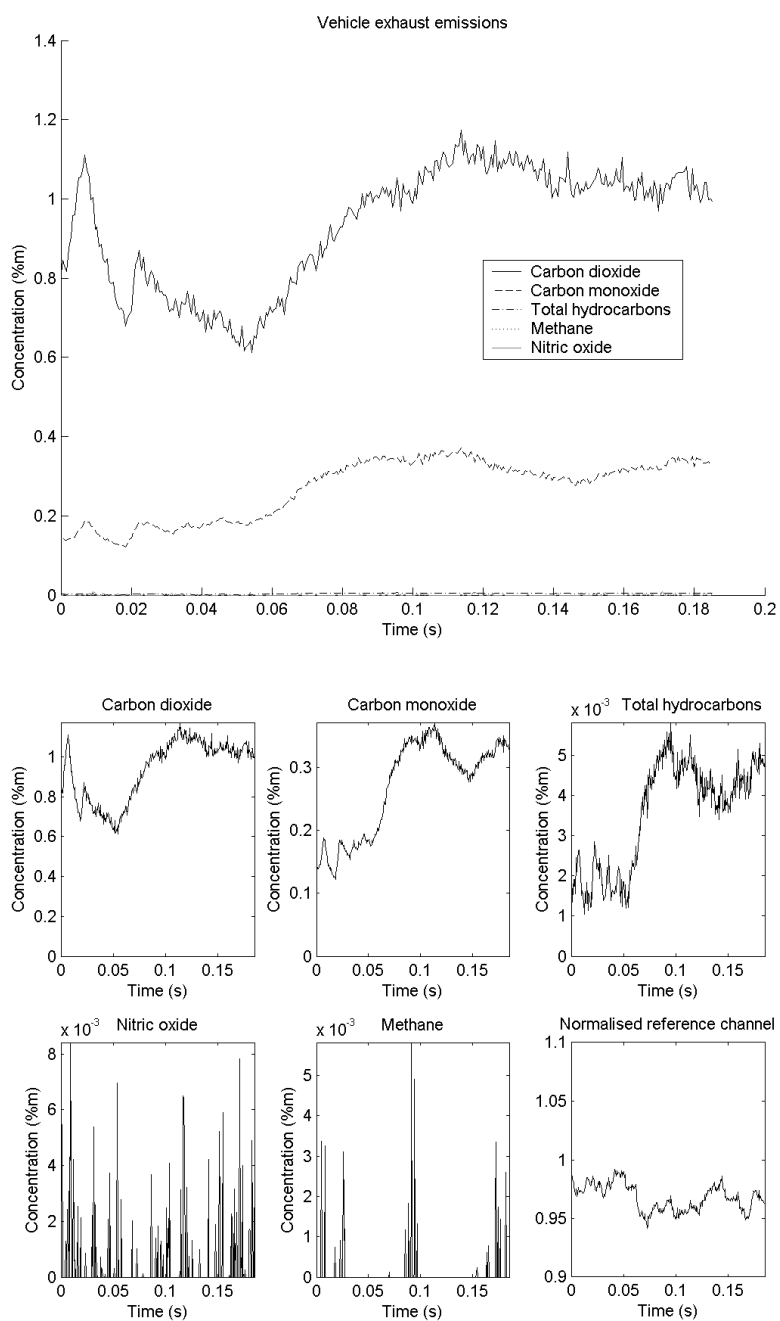


Figure 3.10: Multiple lines and lines which are dashed are best avoided.

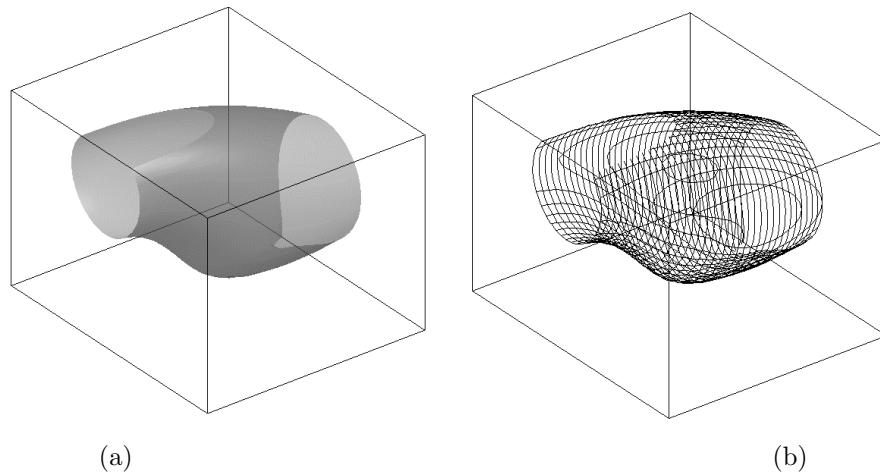


Figure 3.11: Two-dimensional projections of an implicitly defined surface.

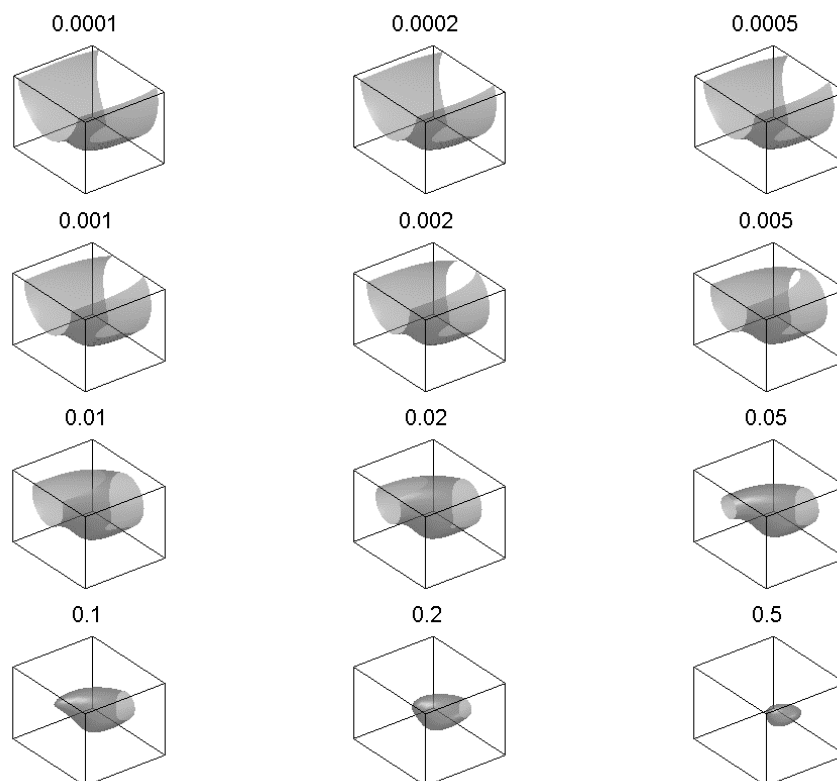


Figure 3.12: Isosurfaces of a volumetric data set $v = f(x, y, z)$.

Volumetric rendering is widely used in the imaging of medical data. The technique produces a two-dimensional image of a three-dimensional data set by calculating some function (such as the maximum value or average value) of the data values along a line, perpendicular to the plane of the image, through the data set and plotting the value at the point where the line intercepts the image. The technique works best with data that contains complex structure and clear boundaries, usually related to a discontinuous jump in function value, between regions (for example between bone and tissue). The technique is less suited to the kind of data sets which are typically encountered in metrology. For example, Figure 3.13 shows a volumetric rendering of the data shown in Figure 3.12. The soft, diffuse nature of the data (caused by the slowly-varying function) makes Figure 3.13 difficult to interpret.

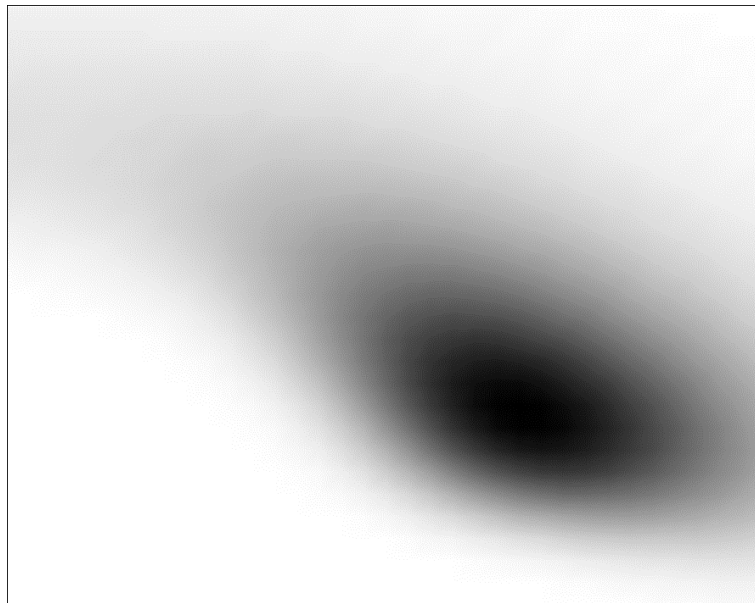


Figure 3.13: A volumetric rendering of the data set displayed in Figure 3.12. The relatively featureless data set makes this image difficult to interpret.

The technique of using a series of plots that are spatially distributed can also be used for four-dimensional data although each individual plot must be very clear for this to work effectively. Figure 3.14 shows a four-dimensional data set $y = f(x_1, x_2, x_3, x_4)$ plotted as a series of isosurfaces. Each plot is a projection of three dimensions x_1 , x_2 and x_3 . The threshold for the isosurface varies from left to right and variable x_4 from top to bottom. It should be noted that variable x_4 is visualised on a different basis to variables x_1 , x_2 and x_3 , and only limited three-dimensional information is available in each sub-plot.

In addition to directly using spatial dimensions as visualisation dimensions to express information in terms of position, other spatial characteristics can also be used, such as length, angle/slope, area, and volume. These are discussed more fully in section 3.2.3.

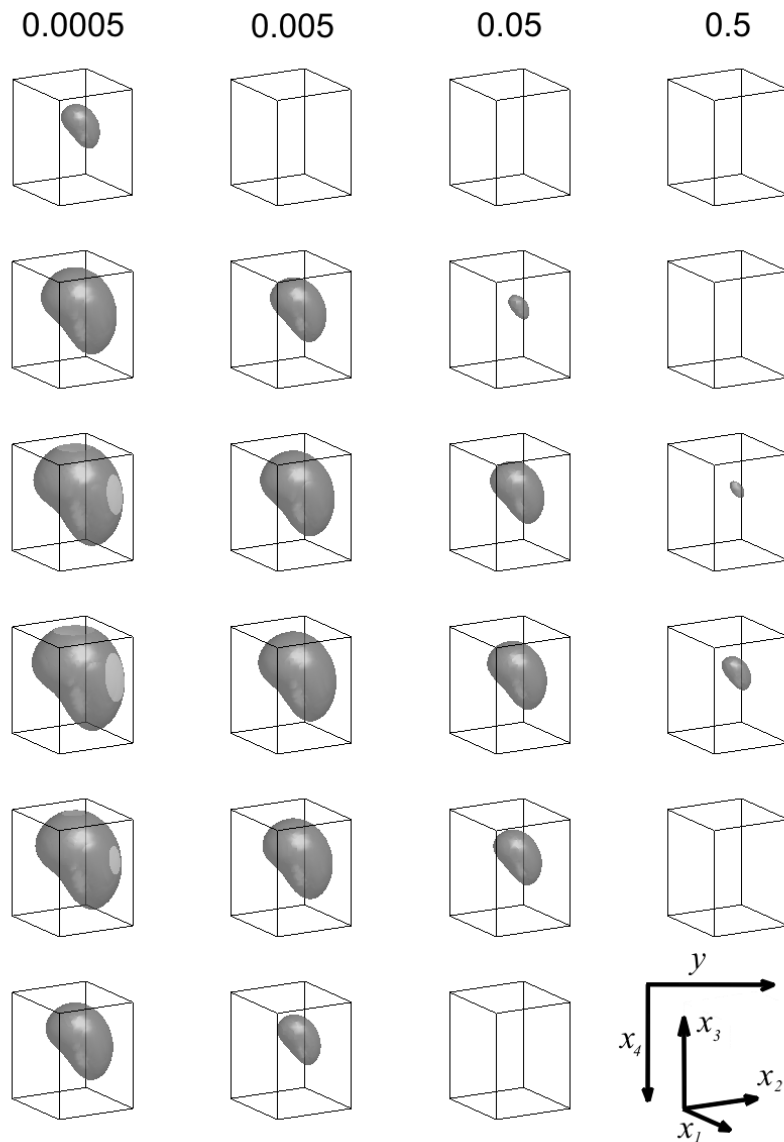


Figure 3.14: Function of four variables visualised as an array of three-dimensional projections of isosurfaces.

Recommendations

- Reserve the use of spatial visualisation dimensions for data dimensions that are important and for those which have similar structure. For example, use spatial visualisation dimensions for data dimensions that have a Euclidean metric, or use angle to represent a data dimension that is cyclic.
- Do not use three-dimensional projections gratuitously.
- Do not use perspective.
- Avoid occluding important data in three-dimensional projections.
- Where possible, use parallax to give sense of depth.
- Use projections of a three-dimensional model where the data is inherently dependent on three spatial dimensions.
- Static projections of three-dimensional point clouds and lines should use drop lines to indicate depth.
- Favour the use of multiple graphs instead of placing multiple lines on a single graph.
- Avoid the use of dashed lines for data.
- Use grids of small plots to visualise three and four dimensional data sets.
- Use isosurfaces in preference to volumetric rendering unless the data has clear regions with distinct boundaries.

3.2.2 Colour

After spatial dimensions, colour is the most widely used visualisation dimension. Colour science is complex, not least because it is dependent on human perception, but in the context of visualisation we can simply consider colour to be a three dimensional quantity – that is, each colour is a point in three-dimensional space.

Colour maps

Although there are a number of choices for the three dimensions (for example, red, green and blue components), we shall find it convenient to work with hue, saturation and value because they are the basic perceptual properties of colour. Hue is the perceptual property which is simplest – it refers to the quantity commonly referred to as ‘colour’ and is plotted around the edge of the traditional artists ‘colour wheel’. Hue is a quantity which is ordered cyclically. Value is the lightness of a colour – two colours which may have different hues are considered to have the same value if neither is lighter than the other. Saturation is the purity or intensity of a hue. As we reduce the saturation of a colour whilst keeping the value constant then it becomes closer to grey. Saturation and value interact – as the value of a fully saturated colour is reduced towards black or increased towards white then its saturation diminishes.

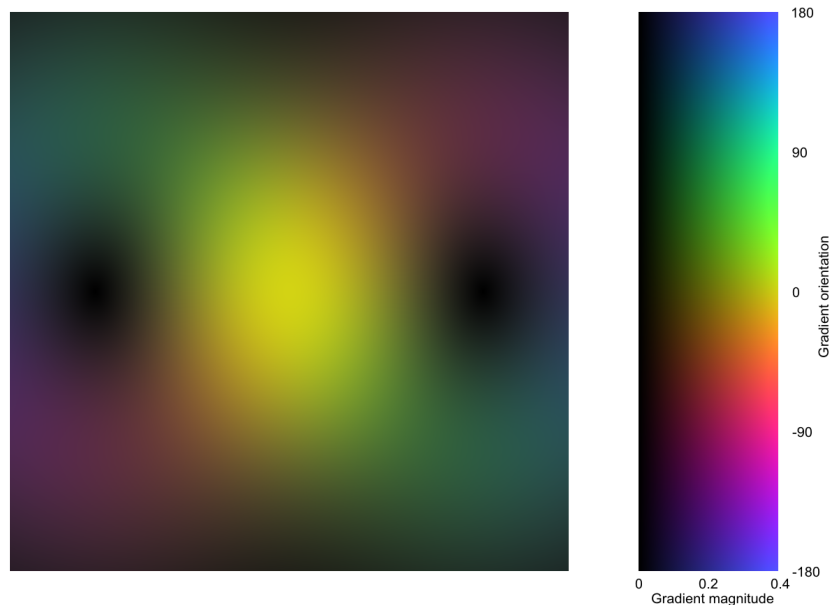


Figure 3.15: Encoding more than one data dimension into visualisation dimensions of colour can be difficult to interpret.

Visualisations should use colour as a single visualisation dimension because it is too difficult to separate more than one piece of information encoded into colour. The rare exception is where colour hue is used to display ordinal data and colour value is used to display a separate characteristic for each ordinal value. Figure 3.15 shows the magnitude and phase of the gradient of the function plotted in Figure 3.8. The data in Figure 3.15 is simple but it is not straightforward to deduce its meaning from the legend. With a more complex data set of the kind that is typically visualised, the use of colour in this way would provide the user with little real information.

Figure 3.16 shows a selection of the colour maps provided with MATLAB. The leftmost three maps encode information into colour value whereas the rightmost four encode information into a single transition between two hues. The ‘pink’ and ‘hot’ colour maps are predominantly encoded in value but a small amount of hue variation is given to enhance the variation in the scale. The ‘hsv’ and ‘cool’ colour maps are encoded into hue – the ‘hsv’ map cycles through the complete spectrum whereas the ‘cool’ map is a limited part of it. Finally ‘jet’ is predominantly encoded into hue with a small amount of value. It is designed to use a large range of hues but not to be cyclic.

With all the colour maps in Figure 3.16, a step in the encoded variable gives a different perceived change in the colour at different points on the scale, and the mapping from colour to perceived magnitude of the variable is quite crude. Normally, the use of colour requires a legend to give quantitative information and this means that users need to keep switching their attention between the plot and the legend.

The flexibility of the mapping between values and colour can be exploited to highlight features within the data. The plots in Figures 3.17(a) and 3.17(b) differ only in the

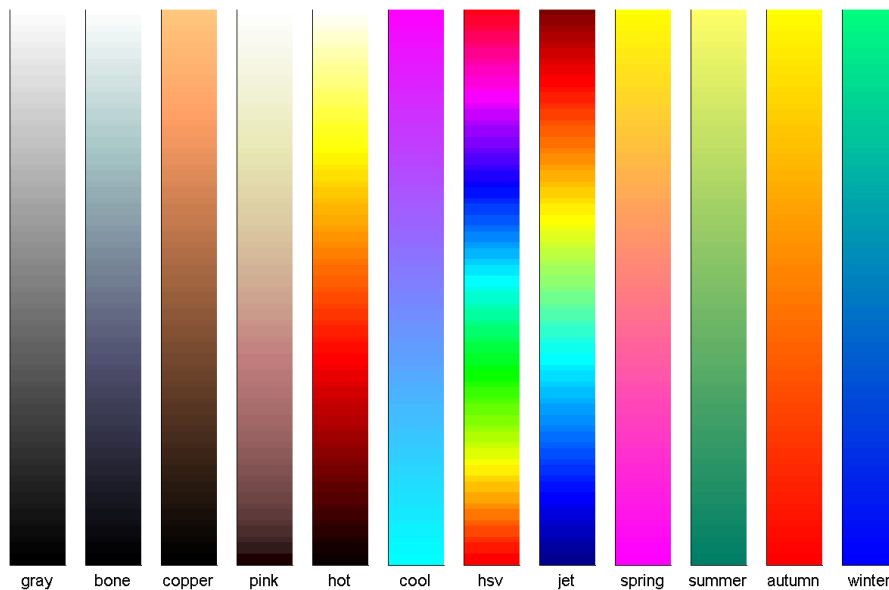


Figure 3.16: A selection of the colour maps provided by MATLAB.

distribution of greyscale values. In Figure 3.17(a) a linear scale has been used whereas in Figure 3.17(b) a power function has been applied in order to concentrate more of the greyscale values in the low range of the plotted variable. For exploratory visualisation, using readily available colour maps (such as a linear greyscale one) is usually sufficient but when designing visualisations for others or when wanting to investigate particular features of the data then the colour map should be designed for the task in hand.

Figures 3.17(c) and 3.17(d) show the use of hue. The human visual system favours colour value variation for high spatial frequency information whereas hue and saturation is more favourable for showing gradual changes in the spatial structure. Figure 3.17(c) highlights a number of problems with a badly chosen colour map. The data gives the impression of being segmented and the linear structures seen within Figure 3.17(a) are missing. In addition, the small localised feature towards the top right of the data is much less obvious.

Issues with colour

Colour has psychological as well as perceptual aspects that should not be ignored when mapping data values to colour. The simple example is mapping a temperature value to colour. Most people would expect the higher temperatures to be mapped to red and the lower values to blue. It can be more difficult with greyscale mappings. Depending on the context it may be more natural to map either black or white to the highest value.

Further factors which should be considered when using colour in visualisations are the normal variation of colour perception in the population and the accurate reproduction of colour. A significant part of the population has some form of colour vision deficiency and this should be considered when designing visualisation. In general, using colour value rather than hue is preferred and where hue is used then reliance on the distinction between

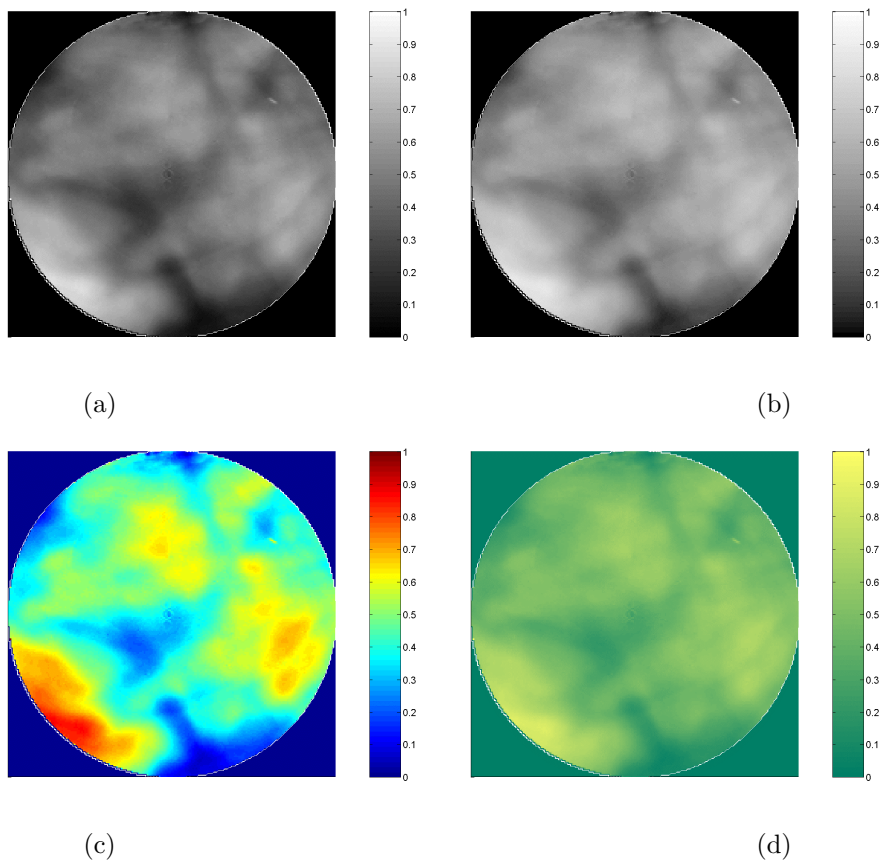


Figure 3.17: The same dataset plotted using four different mappings between variable value and colour.

red and green (which is the most widespread type of colour blindness) should be avoided.

Although good colour printing is now readily available, precise reproduction of colour is difficult and in most visualisation applications it is not possible to accurately predict the colour presented to the user of the visualisation or their perception of it. A colour will appear differently when it is printed as opposed to being displayed on a screen and it is also influenced by the context in which it is placed (lighting and adjacent colours). As such, the interpretation of a visualisation should not rely on subtle distinctions between colours.

The advent of cheap colour and greyscale printing means that the use of cross-hatching to distinguish values now has no place within good graphical presentation of data.

Finally, when using colour in visualisation, the question arises as to whether visualisations should be aesthetically pleasing. The overriding objectives of visualisation are clarity, efficiency, and accuracy – in metrology, visualisations only need to be sufficiently attractive not to detract from those objectives.

Recommendations

- Avoid the unnecessary use of colour.
- Encode only a single variable as colour.
- Use colour value (lightness/darkness) in preference to hue.
- Design the mapping from data magnitude to colour value to be appropriate for the data set and the features which are being highlighted.
- Ensure that people with a colour vision deficiency can use the visualisation.
- Do not rely on fine distinctions between colours to convey useful information.
- Do not use cross-hatching.

3.2.3 Glyphs

Glyphs are discrete objects whose properties are used as visualisation dimensions. Figure 3.18 gives an example of glyphs that take the form of small broken ellipses. The data shown in this figure is the modification of light as it passes through a birefringent material. The input light may have a range of polarisations and the output can have both differing polarisation and absolute phase shift. The shape and orientation of the ellipse show polarisation and the position of the break in the ellipse shows the absolute phase shift. The same glyph is also able to show the attenuation of the light although the data shown is for a material which only affects polarisation. This form of representation is able to display the way in which a five-dimensional quantity varies as a function of two variables. In addition, the glyphs have a natural form that readily relates to the experience of optical engineers. Note that one way of thinking about this representation is as an extension of the technique used in Figure 3.10 – each ellipse can be thought of as a plot of a complex number over time.

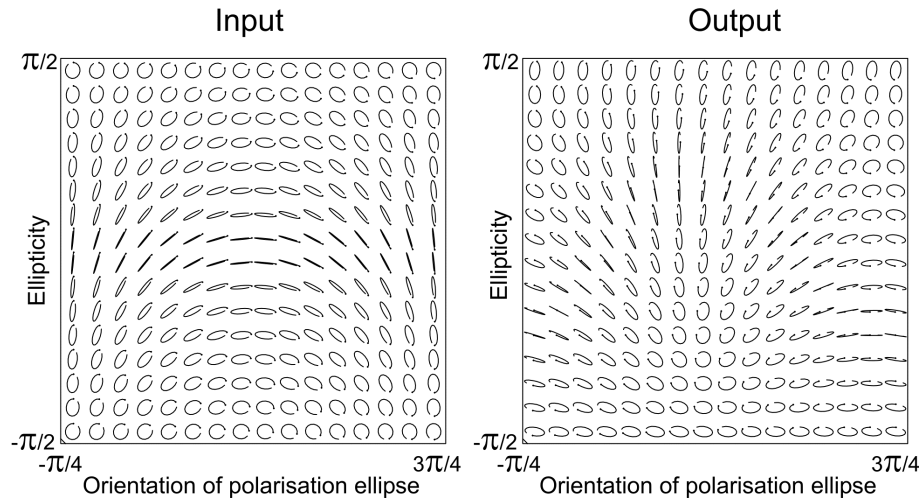


Figure 3.18: The use of glyphs to show the change in polarisation of light as it passes through a birefringent material.

Common glyphs

The simplest glyph is an oriented line giving the direction and magnitude of a vector-valued quantity. These are used to show vector fields – often with an arrow head to indicate vector direction. Figure 3.19 shows a plot of heat flux density data using arrows. The figure shows that the heat flux density is small throughout most of the area shown, but that there is significant flow in the top and bottom right-hand corners. The range of magnitude of the heat flux density gives a large range of arrow lengths. If the large ones do not overlap then the small ones are too small for their direction to be easily seen. Non-linear mapping of the variable values to the arrow length can reduce this problem but at the expense of losing the intuitive connection between arrow length and heat flux density value.

Other commonly used glyphs are flow lines. These indicate the paths of points that always travel in the direction of a vector field. Although they are most commonly used for fluid dynamics data, they can be used for any vector field. Flow lines provide directional information but do not directly provide a sense of the magnitude of the vector field unless it can be inferred from the underlying data (for example, in an incompressible fluid, closer flow lines may indicate higher speed).

Flow line data is not normally calculated for vector fields by the user – instead, it is usual for the visualisation package to calculate them. This can cause problems because artefacts may arise from the algorithms employed (see Chapter 4). Figure 3.20(a) shows a set of flow lines for the thermal data shown in Figure 3.19. Three of the flow lines are seen to coalesce in the region where the thermal gradient drops close to zero. This is a result of the algorithm used to calculate the flow lines not fully reflecting the physical reality of the experiment.

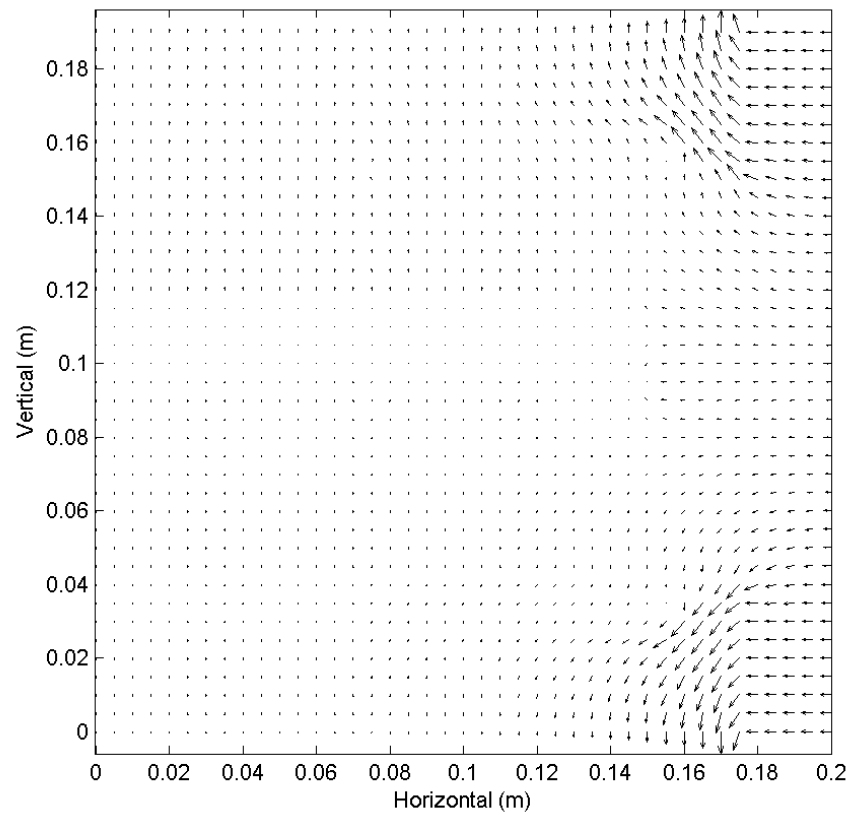


Figure 3.19: Heat flux density displayed using arrows.

Flow lines can also be used for three-dimensional vector fields, as in Figure 3.20(b). To make such projections understandable it is usually necessary to provide depth cues by thickening the lines and using lighting but, without the ability to rotate the view point around the model, these cues give a poor appreciation of the three-dimensional structure unless that structure is very simple. Figure 3.20(b) also shows that the cross-section of the flow tube can be used to display the variation of a further variable along the flow line. The typical use of these dimensions is to show the divergence of a vector field in the overall size of the cross-section and to show its curl by flattening the tube to a ribbon and rotating the cross-section of the ribbon about the flow-line.

Visualisations of other data which varies along the flow line – or, indeed, any path in three-dimensional space – can be constructed by further elaboration of these ideas. Thus, Figure 3.21 shows how several different data sets which vary along a path can be visualised. As in Figure 3.20(b), we show the path using a generalised ribbon which has a finite cross-section, and also a cross-sectional shape. Then one variable can be mapped to the colour of the ribbon (as in Figure 3.21), while another to its cross sectional area and a third to its rotation (as in Figure 3.20(b)). We can add glyphs along the path to show the variation of more variables – in Figure 3.21, one is mapped to the colour of the glyph, but others could be mapped to its size, or (if the glyph is anisotropic) to its orientation.

The use of glyphs to represent uncertainty in two-dimensional plots is commonplace. Error bars and box plots (Figure 3.22(a)) enable more information to be displayed than just the measured variable or the mean of a set of values. Displaying the covariance matrix for a pair of variables as an ellipse is also well-established. This can be readily extended into three dimensions by the use of ellipsoids (Figure 3.22(b)). Such ellipsoids can be thought of as an isosurface of the Gaussian (normal) probability density function with the given covariance matrix. Where the probability density function is known, its isosurface can be used. The data shown in Figure 3.12 represent a four-dimensional probability density function.

Issues with glyphs

Glyphs are discrete objects so their display requires them to be spatially separated, usually on a fairly sparse grid. Glyphs also have finite size. Both of these characteristics lead to the potential for misrepresenting data by under-sampling and poor localisation of the value. In Figure 3.19, only the regular starting points of the arrows indicate that the value is localised at the start of the arrow rather than its mid point. The localisation convention used should be clear and the spatial extent of the glyph should be no larger than the region in the underlying data where the values are roughly constant.

While the ellipses in Figure 3.18 and the arrows in Figure 3.19 have a linear mapping between physical size and perceived size, this is not the case for all glyphs. Area is more difficult to judge than length and it may not be clear whether the data dimension has been mapped to the linear or area dimension of the glyph. Adding a legend to a diagram can help to specify the dimensional mapping. However, few visualisation packages make this facility easily available despite its common use in geographic information visualisation.

While a wide variety of glyph forms can be devised, it is important that they are obvious to the user of the visualisation. Attempts have been made to encode stock-market data as characteristics of architectural forms with the belief that the mapping from data to

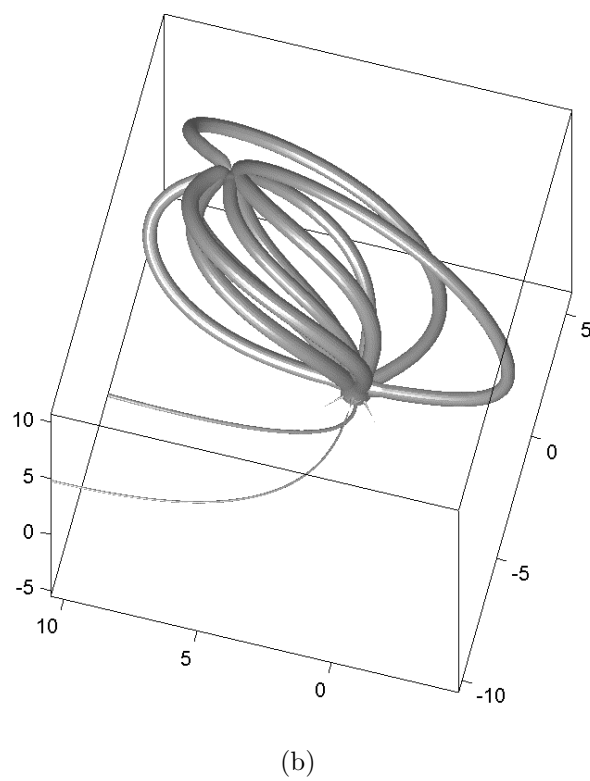
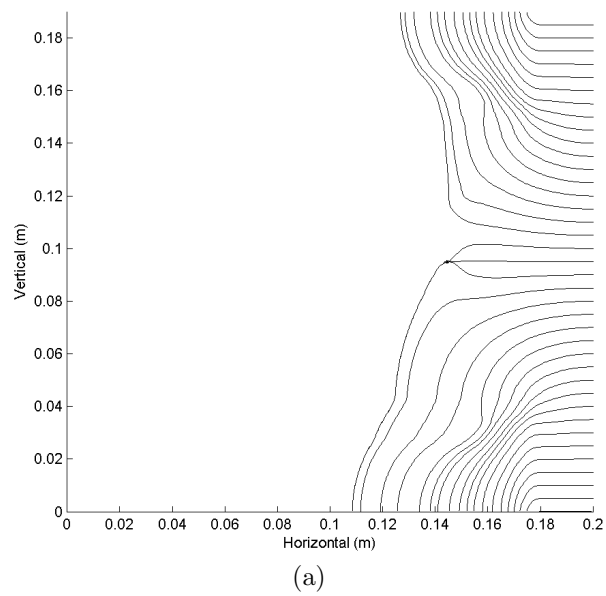


Figure 3.20: Stream lines and stream tubes can be used to show flow lines in two and three dimensions.

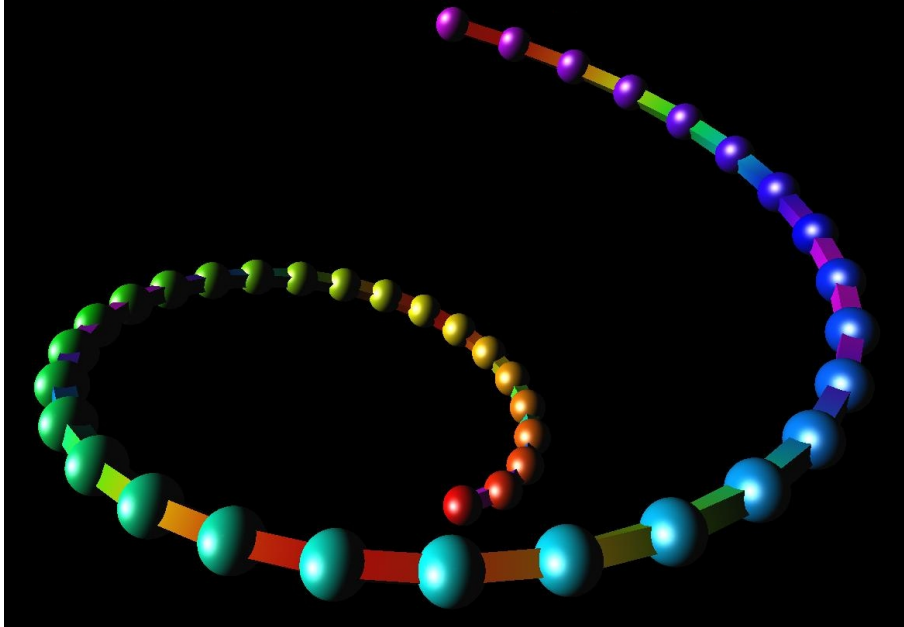


Figure 3.21: Displaying the variation of several values along a path in three dimensions with a generalised ribbon and mapping data values using colour, glyphs and other properties.

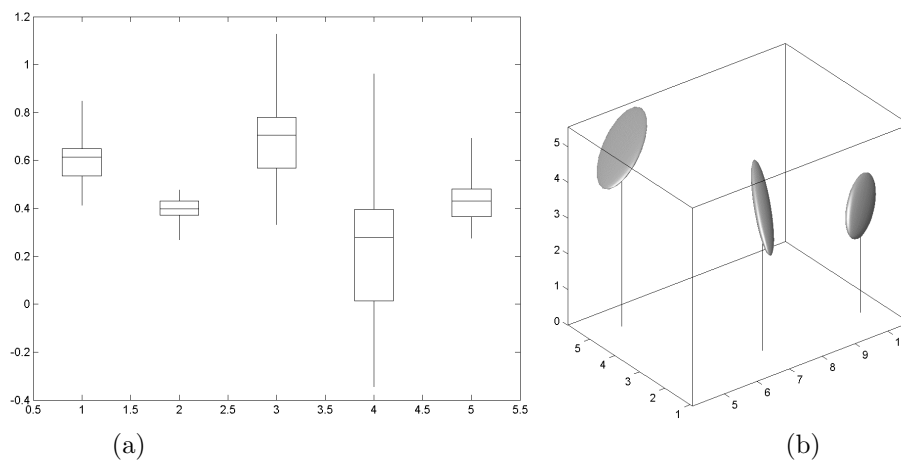


Figure 3.22: Representing uncertainty: (a) standard box plot, (b) 3x3 covariance matrices represented as ellipsoids.

visualisation could be constructed so that normal data would correspond to an aesthetically pleasing building and abnormal data to an ugly building. Such forms of visualisation are unnecessarily obtuse and take time for the user to learn how to use. In the visualisation of metrology data, glyphs should be simple so that the user can readily understand them.

Recommendations

- Use simple to understand glyphs where there is a natural mapping from the data to the properties of the glyph.
- The perceived position of the glyph should coincide with the data it is representing.
- The spatial extent of a glyph should be no larger than the region of the underlying data that can reasonably be represented by the glyph.
- Use legends to give a quantitative understanding of glyph value.

3.2.4 Displaying edge- and face-based data

The discussions so far have focussed on data values which are associated with points in space. Some types of data, however, may be such that the data is associated with objects of a higher dimensionality than points – for example, lines (or edges), areas (or faces) or volumes (or cells). Consider flow through a network which consists of straight channels connecting nodes. The channels are one-dimensional objects, whilst the nodes have zero dimension. Such a network could be used to model fluid flow through a porous solid, or the flow of traffic along roads which connect towns together. The flow along a channel is assumed to be constant along its length, but it is a property of the channel, not of the nodes. Thus, to visualise the flow, it must be somehow associated with the representation of the channel, using (say) a colour. See Figure 3.23(a) for an illustration.

Again, consider a finite-element model of the properties of a machine part. The geometry of the part is modelled using a three-dimensional mesh in which vertices are connected by edges, edges make faces and faces make cells. Some of the data calculated in the finite-element model is associated with the vertices (temperature, for example), whilst other values are calculated as properties of (say) faces – for example, the stress on a face, or the area of a face. One can approximate these variables before visualisation or analysis by arbitrarily dividing them up and assigning them to neighbouring vertices, but sometimes it is useful to display the data exactly as the calculation has produced it, and here, techniques for data visualisation associated with objects of dimensionality higher than zero are required. Figure 3.23(b) shows an example where temperature has been calculated for each cell, so each cell is coloured according to the temperature within it.

Visualisations of data associated with higher-dimensional elements can also make use of culling to show regions of interest. Thus, in Figure 3.23(a), it would be easy to imagine only displaying those channels for which the flow was in a given range. Or, in Figure 3.23(b), removing the cells whose temperature was less than some value. Such treatments are not readily available for data associated with points, since points cannot be easily removed from a three-dimensional structure.

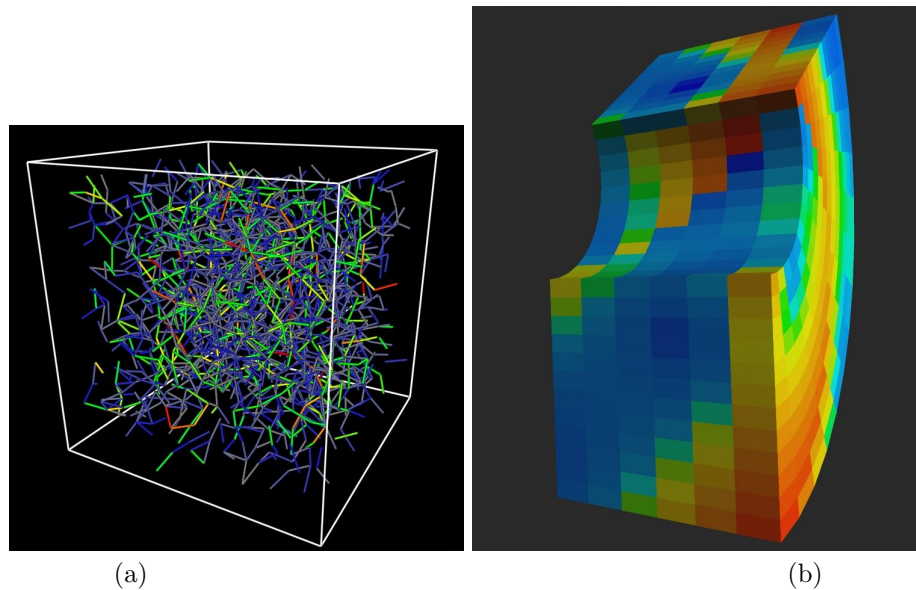


Figure 3.23: Displaying data associated with dimensions higher than zero. (a) Flow through a network of straight channels connecting nodes. (b) Temperature in a machine part, calculated using finite-element modelling.

Recommendations

- Map visualisation data to objects of the appropriate dimensionality (points, edges, faces, volumes).
- For data associated with dimensions higher than zero, make use of culling to drill into a visualisation to reveal interior structure.

3.2.5 Animations and interactivity

Time can be used as a visualisation dimension by changing the parameters of the visualisation at each time step. At its simplest this gives a passive animation – one that simply runs end-to-end, but enabling the user to vary the time point interactively makes it much easier to concentrate attention on a particular event. Such animations can be particularly effective where time is one of the data dimensions. A classic example of this use of animation was the confirmation by British Telecom of the link between line faults and thunderstorms. Here, they plotted the location of line faults spatially and animated their occurrence in time. Although line faults have many causes, a significant number were seen to follow the path of thunderstorms as they developed. This example exploits the fact that people are very good at perceiving spatial pattern in movement. Some examples of passive animations concerning the visualisation of uncertainty are provided in Chapter 5.

The primary drawback to using time as a visualisation dimension is the limited formats that the visualisation can be delivered in. Although flip-book animation has been used for scientific visualisation, this is a novelty and for practical purposes time cannot be used for

visualisations which are printed. A less important drawback is that time is not readily interchangeable with other visualisation dimensions. In this respect it shares the same characteristics of using multiple plots as in Figure 3.12. Although four of the data dimensions in Figure 3.12 are all equivalent, one gains a different impression of the dimension that spans across the plots rather than being wholly contained within each one. Time used as a visualisation dimension suffers the same drawback.

It is often helpful to allow the user to change parameters of the visualisation interactively over time. A common instance of this form of interaction is allowing the user to alter the viewpoint of a three-dimensional scene, although any visualisation parameter could be varied. In principle, any number of parameters can be explored in this way. However, the technique is only effective if the user can discern a clear pattern between changes in the parameters. With too many parameters, it is easy to become lost in a multiple dimensional space without perceiving any clear structure to the data.

Allowing the user to change the viewpoint of a fixed scene and building animation sequences frame by frame are generally available in good visualisation tools. The ability to construct more general interactive forms is less accessible, and such visualisations can take a significant time to construct. Hence their usefulness is currently limited, although this is expected to change as visualisation tools develop.

Recommendations

- Design time-based visualisations to take advantage of the human ability to perceive spatial pattern in movement.
- Use interactive visualisations to enforce the sense of depth in projections of three-dimensional scenes.
- Restrict the number of variables that can be varied within an interactive visualisation so that users do not become lost within the data set.

3.2.6 Parallel co-ordinates

As was mentioned in section 3.2.1, visualisation dimensions are commonly mapped onto spatial dimensions. This approach is ideal for data sets involving three or fewer visualisation dimensions, and the use of colour can make visualisation of data with four dimensions possible, but as the number of visualisation dimensions rises so does the difficulty of interpreting the subsequent image.

One way of avoiding this difficulty is to use the parallel co-ordinate technique, developed in the 1980s by Professor Alfred Inselberg [18]. A standard spatial visualisation of a data set uses a set of perpendicular axes with a common origin, and often uses equal scaling along the axes so that visualisation dimensions are easily comparable. The parallel co-ordinate technique uses parallel axes with no common origin and (in general) differing scales. The use of parallel axes means that any number of dimensions can be visualised simultaneously. The technique has been successfully applied to process control problems involving several hundred variables [7].

Suppose that we wish to explore a data set $\{\mathbf{x}_i, i = 1, 2, \dots, n\}$ where the $\mathbf{x}_i$ consist of N quantities, so that $\mathbf{x}_i = \{x_i^j, j = 1, 2, \dots, N\}$. The simplest form of parallel co-ordinate plot for an N -dimensional data set consists of a set of N equally spaced parallel axes of the same length. Each axis represents a different data dimension, and is (initially) linearly scaled between the maximum and minimum value of the corresponding quantity within the data set. In some cases it is helpful to apply identical limits to all axes representing the same type of quantity, or to use logarithmic scales, and a good software implementation of the technique would allow for this. Each point $\mathbf{x}_i$ within the data set is represented by a set of straight line segments linking adjacent axes. On the assumption that the axes are parameterised between 0 and 1, the set of line segments links

$$\hat{x}_i^j = \frac{x_i^j - x_{\min}^j}{x_{\max}^j - x_{\min}^j}, j = 1, 2, \dots, N, \quad (3.1)$$

for each value of i , where $x_{\min}^j = \min_{i=1,2,\dots,N}\{x_i^j\}$ and $x_{\max}^j = \max_{i=1,2,\dots,N}\{x_i^j\}$.

As an example, suppose we want to visualise the following five six-dimensional points:

$$\begin{aligned} \mathbf{x}_1 &= (80, 97, 86, 13, 94, 86), \\ \mathbf{x}_2 &= (3, 48, 57, 52, 78, 56), \\ \mathbf{x}_3 &= (43, 67, 26, 85, 97, 94), \\ \mathbf{x}_4 &= (89, 74, 79, 13, 72, 87), \\ \mathbf{x}_5 &= (32, 66, 94, 76, 25, 7). \end{aligned}$$

First the points would be normalised using the formula 3.1 given above, which would give the new points as

$$\begin{aligned} \hat{\mathbf{x}}_1 &= (0.90, 1.00, 0.88, 0.00, 0.96, 0.91), \\ \hat{\mathbf{x}}_2 &= (0.00, 0.00, 0.46, 0.54, 0.74, 0.56), \\ \hat{\mathbf{x}}_3 &= (0.47, 0.39, 0.00, 1.00, 1.00, 1.00), \\ \hat{\mathbf{x}}_4 &= (1.00, 0.53, 0.78, 0.00, 0.65, 0.92), \\ \hat{\mathbf{x}}_5 &= (0.34, 0.37, 1.00, 0.88, 0.00, 0.00), \end{aligned}$$

to two decimal places. Most parallel co-ordinate software packages will carry out the normalisation automatically, but if a non-specialist package is being used then the user will need to normalise the data before plotting. Finally, the data points are represented as line segments linking the parallel axes. The visualisation of this data set is shown in Figure 3.24. Reading the axis representing Q1, the first quantity, from top to bottom, the line segments represent $\mathbf{x}_4$, $\mathbf{x}_1$, $\mathbf{x}_3$, $\mathbf{x}_5$, and $\mathbf{x}_2$.

The line segments representing each point are easy to identify in this example. For instance, the line segments that link the top end of the final three axes must represent $\mathbf{x}_3$ because the final three normalised co-ordinates of $\mathbf{x}_3$ are 1.00. Identification of individual points becomes more difficult when more points are plotted, but some software packages

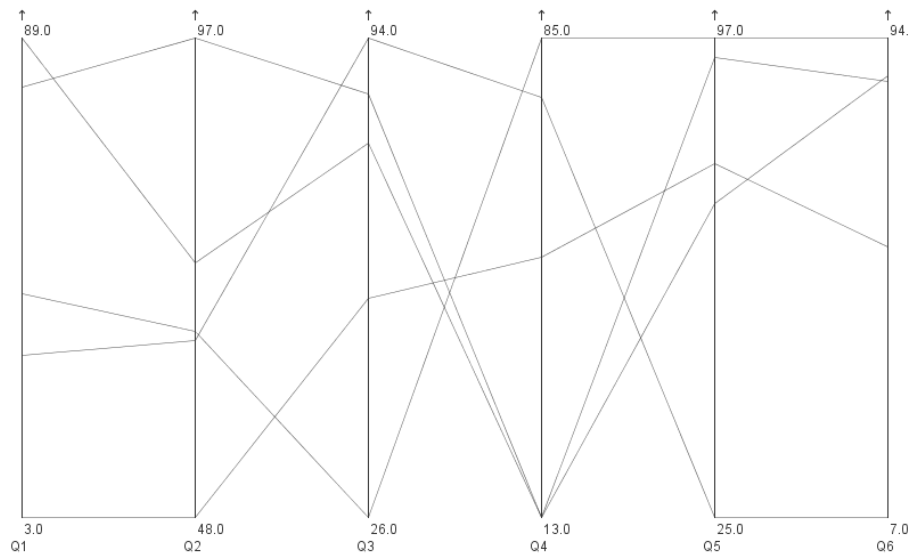


Figure 3.24: A simple example of a parallel coordinate plot.

have tools to overcome this difficulty by allowing the user to highlight subsets of the data interactively.

The above example is a trivial one, and does not illustrate the best use of the parallel co-ordinate technique since it involved few points and few dimensions. Parallel co-ordinate visualisations are most useful for large data sets of high dimension. In addition, the technique is most useful when the images are interactive. The technique is very good for exploring data rather than for presenting static images to be interpreted by others. Interactive visualisations of this type enable the user to focus on subsets of the data in order to explore the relationships between different visualisation dimensions. The interactivity requirement means that whilst packages such as Excel and Matlab can be used to create parallel co-ordinate plots, they do not easily enable the user to explore the full potential of the technique.

Plotting large data sets as parallel co-ordinate plots usually leads to a large number of intersecting line segments between each adjacent pair of axes. The outer limits of these intersecting line segments sometimes forms a smooth curve, and the shape of this curve can indicate the form of a relationship between the quantities corresponding to the axes. As an example, the upper limit of the highlighted lines between the first two axes in figure 5.6 forms a smooth curve typical of points lying within a circle or annulus.

An introduction to the parallel co-ordinate technique and its application to measurement data has been prepared as part of the current SSfM programme [42]. The report includes a more extensive discussion of the strengths and weaknesses of the technique, as well as two case studies of its application to real data.

Recommendations

- Parallel co-ordinate plots can provide a good way of looking at large data sets of high dimension.
- The technique is best used as a way of interactively exploring data, rather than as a way of creating static images.
- Simple versions of the technique can be implemented using Excel or Matlab, but these versions lack the interactivity and special exploration tools that more specialised packages include.
- See the report [42] for more details on the technique.

3.2.7 Combining visualisation dimensions

Sections 3.2.1 to 3.2.6 have described visualisation dimensions and techniques in isolation, but part of the power of visualisation arises from combining the techniques together. A correct balance has to be made between encoding information into a single design where all the data is accessible in a single image and the use of multiple images. It is strongly recommended that when visualising data a large number of simple plots are first prepared. These are generally quick to produce and can show the overall features that are most important. Once this overview has been obtained more complex visualisations can be built up in stages.

Figure 3.25 shows colour and glyphs used in combination to visualise both the temperature (a scalar-valued quantity) and heat flux density (a vector-valued quantity) within an experimental set-up. Also plotted in the figure are the material boundaries so the temperature and flux density can be related to the physical apparatus used for the experiment. The heat flux density data of this figure is the same as that displayed in Figure 3.19. When constructing plots such as this, it is necessary to ensure that each visualisation dimension is clearly visible – a balance must be made between collocating information on a single plot for easy comparison and making that plot too complex to understand. Plotting data as both separate and combined plots can be particularly effective for creating understandable visualisations.

When mapping several data dimensions to spatial visualisation dimensions, it is necessary to consider the types of the data dimensions and the relative scales of their mapping to the visualisation dimension. If the data has a single type and there is no scaling of the values, then it is straightforward for the user to understand the values within the visualisation. An example of this situation is the box plots shown in Figure 3.22(a) where the dimensions of the glyphs are on the same scale as the underlying data. Scaling one of the variables can introduce misleading artefacts into the visualisation. Figure 3.26 shows the error in the radius of a perfect sphere mapped to the same dimension as the sphere. The actual departure from a perfect sphere is very small and would be negligible if plotted on the same scale as the sphere. Figure 3.26 does not show the true surface form of the sphere, although provided the relative scaling is apparent to the user of the visualisation then this plot can be very informative. If the different visualisation dimensions are perceived in different ways, then relative scaling does not cause perceptual problems. This is the case in Figure 3.25: the colour scale used to visualise temperature and the glyphs used to visualise

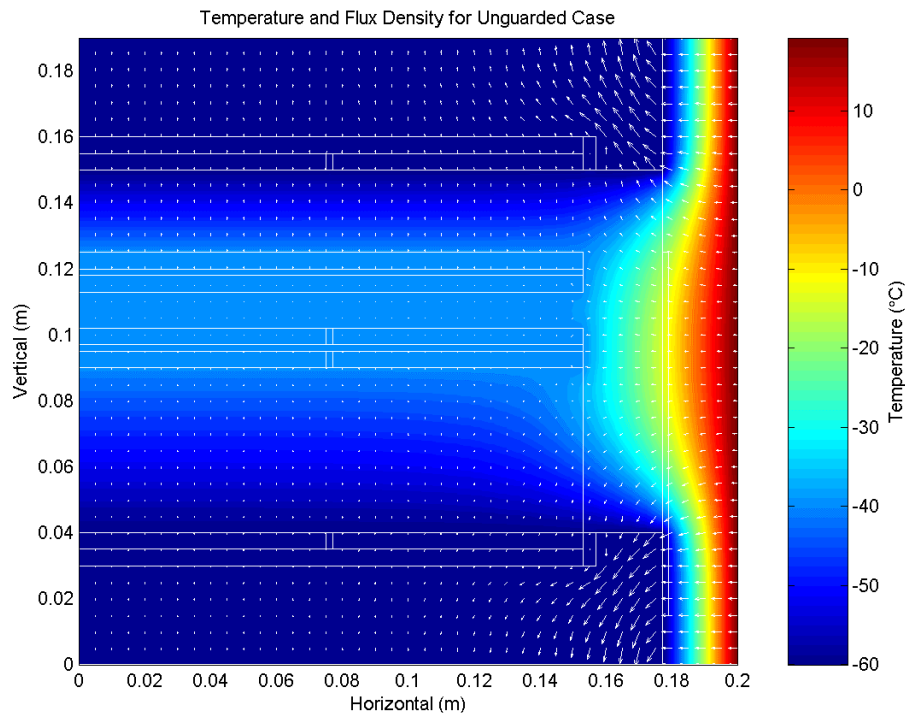


Figure 3.25: Thermal data mapped to both colour and directed arrows.

heat flux density are perceived separately and differently, so the different scales that are used to map from the data dimensions to the visualisation dimensions do not affect the observer's understanding.

A wide variety of plots can be combined. One objective in using a combination of plots is to allow each one to be seen in the context of the others. The simplest form of combination involves the use of positioning and identical scales (Figures 3.10 and 3.14). In addition, slices through data can be combined with projections of three-dimensional models in order to show values on particular planes. A variant of this technique is to use animation to sweep the plane through the data set or to allow its position to be adjusted interactively. Figure 3.27(a) shows a slice through a volumetric data set combined with an isosurface. This approach can place the slice within the context of other data but can be difficult to evaluate. Combining the projection of the three-dimensional model with a separate plot of the slice data, (b), can make quantitative assessment of the data easier.

The notion of taking a slice through data to produce a lower dimensional plot can be extended to other dimensions. We have already seen (Figure 3.14) how a four-dimensional data set can be displayed as a collection of three-dimensional plots by taking slices (or projections) in three dimensions. Figure 3.28 illustrates slicing through a two-dimensional data set, using a one-dimensional slice. Here, the slice is parallel to one of the two dimensions (x or y) of the plot. The slice in (a) is also colour-mapped by data value: the colours run (low values to high) from red through yellow and green to cyan. The plot and the colours on the slice update as it is moved through the data set. Figure 3.29 shows the effect of taking a one-dimensional slice through a three-dimensional data set – here, the

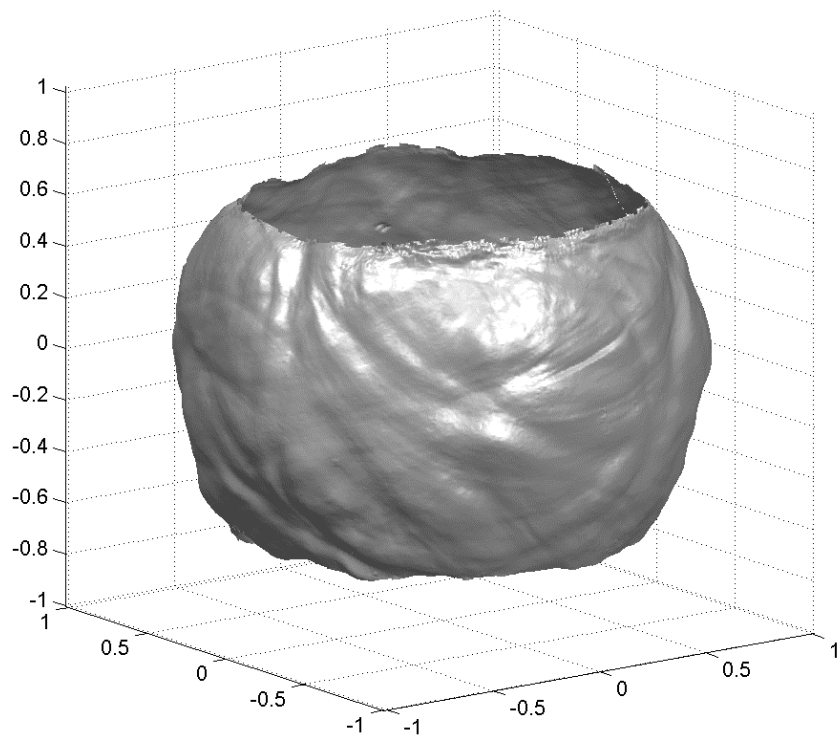


Figure 3.26: Departure from a perfect spherical surface plotted as a function of spatial position but greatly exaggerated in scale.

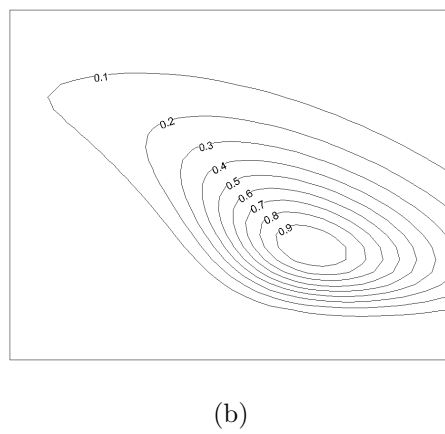
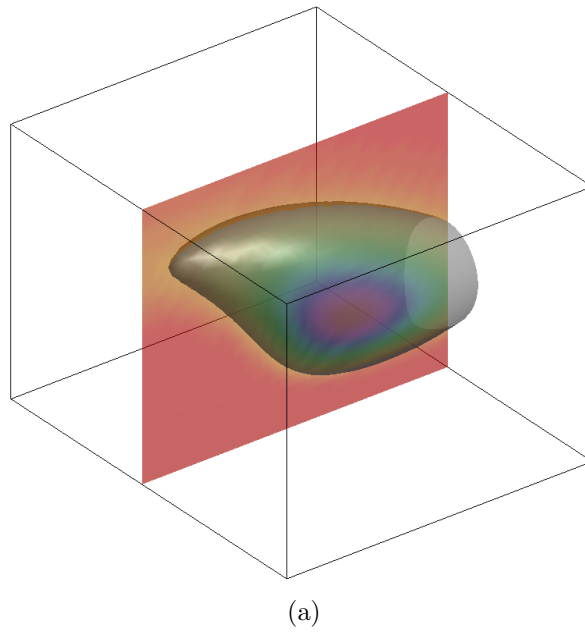


Figure 3.27: A slice from a volumetric data set plotted (a) with an isosurface, and (b) as a separate plot.

location, orientation and the length of the probe are all under user control. The data set (here, the potential energy between three atoms) is displayed using an isosurface and the one-dimensional probe is represented as an arrow. Figure 3.29(b) shows the way the data values (potential energy) vary along the length of the probe from the tail to the head of the arrow. The plot updates as the probe is manipulated in the three-dimensional visualisation. Note the usage of rendering effects such as transparency, highlighting and shading in Figure 3.29(a) enhance the appearance of the isosurface.

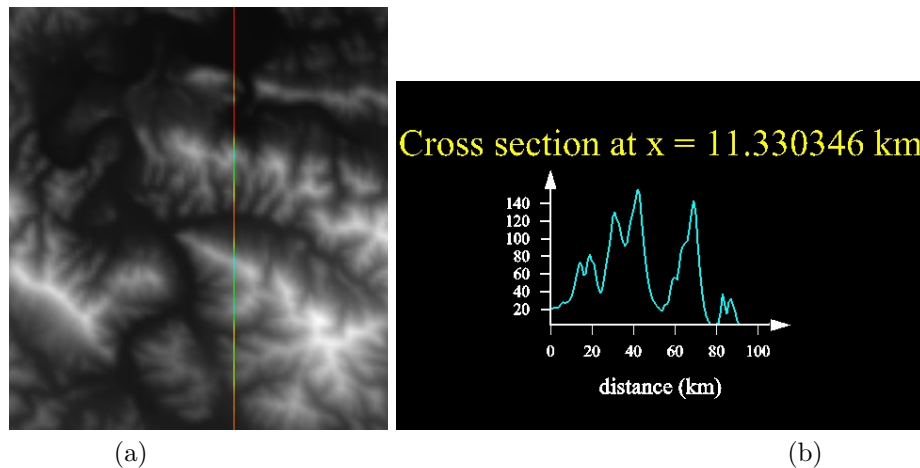


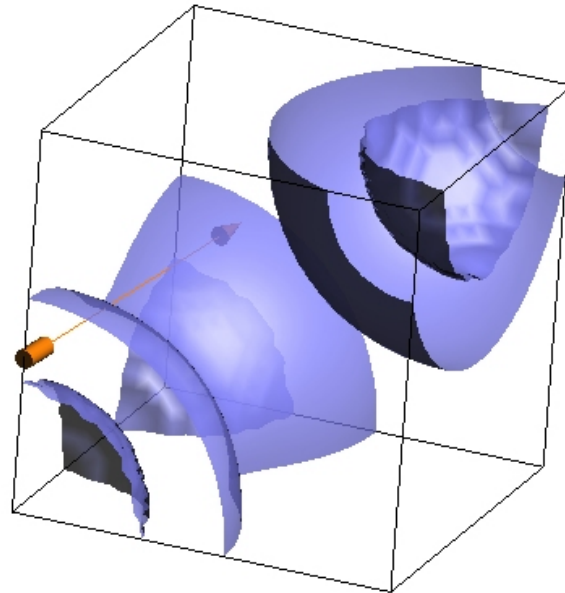
Figure 3.28: Slicing through a two-dimensional data set – here, a height field. (a) Data set displayed as a greyscale contour map, with larger values mapped to white. The coloured line indicates the location of the (one-dimensional) slice. (b) Plot showing how data varies along the slice. The origin of the plot is at the bottom of the contour map (i.e., small y).

Recommendations

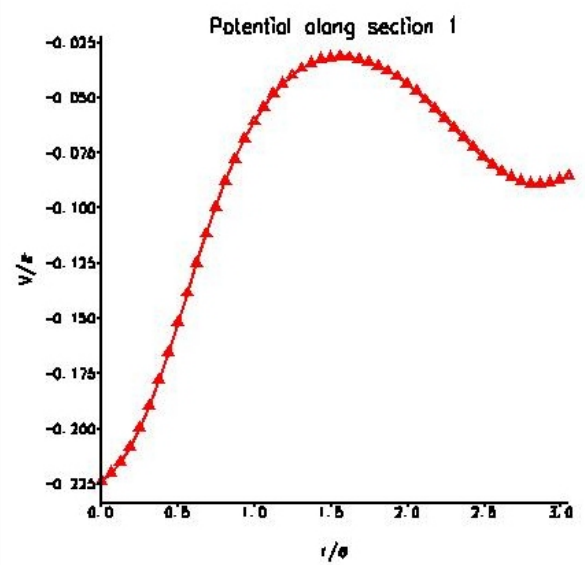
- Avoid unnecessarily complex visualisations by using multiple plots.
- Balance the visual impact of different sources of information to ensure that they are all clearly visible.
- Use separate plots in conjunction with combined ones where it leads to clarity.
- When mapping several data dimensions to several spatial dimensions using different relative scalings for each map, consider whether a misleading impression of the data is being given.

3.3 Mapping the quantity being visualised to a visualisation dimension

Many data sets comprise just simple relationships between pairs of variables. This kind of data is relatively easy to plot and the process of visualising it is familiar to metrologists.



(a)



(b)

Figure 3.29: Slicing through a three-dimensional data set (a) using a one-dimensional probe to produce a line plot (b).

The difficulties arise when there are many more variables in the data and when the underlying structure of the data is not known.

It is a common misconception of people unfamiliar with visualisation that visualisation will permit them to display simultaneously an arbitrary number of dimensions, whereas in practice the number of visualisation dimensions that can be reasonably perceived is limited. Conversely, it is also incorrectly assumed that visualisation can only be applied to very low dimensional data. The creative aspect of data visualisation is designing effective mapping from the data to the visualisation.

If the structure of the data set is known then it is desirable to map to visualisation dimensions that preserve that structure. For example, if the normal Euclidean distance metric holds, then it is appropriate to map the data to equally scaled spatial axes. Similarly cyclic quantities map well to colour hue and to angle as both of those visualisation dimensions can be configured to be cyclic. As with all data analysis, in visualisation it is always sensible to carefully consider the nature of the data.

Often the structure of the data may not be expressed in its original form and it is worth considering transforming the variables so that the structure becomes apparent. Depending on the nature of the data, a transformation between Cartesian and polar co-ordinates may be appropriate or it may be useful to map the ratio of variables to the visualisation dimension. Identifying symmetries in the data set can also lead to a reduction in the dimensionality of the data set being visualised. Such techniques are highly dependent on the particular data set and their applicability may only become apparent after the initial stages of exploring the data set with visualisation.

With high-dimensional data sets where the structure is not known, techniques such as Principal Components Analysis (PCA) can give derived variables that provide greater insight than using the given data variables. Even with inherently non-linear data sets PCA has been proven to be effective, for example, it is one of the standard approaches for summarising the variation in images of peoples' faces. Although techniques such as PCA can give a reduction in dimensionality, it is recommended that they are never applied blindly – interpretation of data can never be completely automated. Normally techniques such as PCA provide a starting point for visualisation and once the structure of the data set starts to become apparent then that structure can be directly exploited.

When mapping data to a visualisation dimension, the transformation of the data value has to be considered. For example, if we wish to map a variable to a linear spatial dimension then options include simply using a scaled version of the variable or using a logarithmic function. The mapping used should be designed to emphasise the important characteristics of the data. When choosing the mapping it is important to consider the perceptual scale of the visualisation dimension and to map the important data values to the most easily perceived part of the scale of the visualisation dimension. Figure 3.17 demonstrates that not all mappings are perceived identically.

As well as considering the perceptual scale of each visualisation dimension, it is also necessary to ensure a reasonable balance in the perception of different visualisation dimensions. In Figure 3.14 the perception of the x_4 variable differs from that of the other variables. If x_4 is no different from the other variables from the point of view of the underlying data set, then it may be desirable to create additional versions of the plot with x_4 interchanged with each of the other variables in turn.

While good design and clear reasoning should always be the basis for mapping data to visualisation dimensions, it is best to explore a number of such mappings in order to identify the best ones. In exploratory data analysis, multiple views of data are always useful, and when using visualisation to present information, multiple views ensure that the most concise and efficient presentation is used.

Recommendations

- Exploit the underlying structure of the data set to reduce the number of dimensions in the visualisation.
- Preserve the important structure within data sets when mapping data dimensions to visualisation dimensions.
- Take account of the perceptual scale when mapping variables to visualisation dimensions.
- Where it is necessary to use an asymmetric visualisation of data, create alternative forms where the data variables are interchanged.
- Produce multiple views of data sets.

Chapter 4

Implementation of the Visualisation

Misleading visualisations can be created by bad design, as has been explained in the previous Chapter. The use of software to create visualisations means that misleading visualisations [6] can also be created if the visualisation software is inadequate. Errors in visualisation software can easily be missed [17]. It is therefore essential that the software that produces visualisations be validated and tested, and that the results of the validation and testing activities be documented and made available to the users of visualisation software. This chapter suggests some points to consider when implementing visualisations, and some ways of testing and validating them once they have been created.

4.1 Development of visualisation software

Two types of generic software problem exist:

- The visualisation software does not implement the programmer's intention.
- The software performs differently from the expectations of the scientist using it [31].

The development of visualisation software should thus be approached in the same manner as any other software development project, with the preparation of user and functional specifications and the definition of a comprehensive test plan to address the above problems.

If it is necessary to replicate the output and have full traceability of the end form of the visualisation, then the data files, scripts, spreadsheets and other files used to create aspects of the visualisation should be documented and controlled. It is very common when interactively exploring a data set to continually modify scripts and other software as successive plots are produced. Where that exploration is simply to gain an overview of a data set and the graphical output is discarded, full software control is not necessary.

However, if the exploration subsequently becomes a visualisation that others will use and they need to be confident of its accuracy, or if a decision is to be made based on the visualisation, then more formal documentation and version control is required.

The specification of the visualisation software needs to take into account the *purpose* of the visualisation. Is it to give a qualitative understanding of some phenomenon of interest or is the aim to display a definitive representation of an experimental or theoretical result, perhaps including the uncertainties associated with that result? In the context of visualisation research, little attention seems to have been paid to the question of visualising uncertainties [23].

Visualisation is not a passive process – it modifies the data in the course of mapping it to the graphical form. The algorithms used for transformation, interpolation and rendering should be evaluated for their suitability and accuracy. Inaccurate and unsuitable algorithms lead to inaccurate and misleading visualisations. Additionally, the choice of visualisation method may lead to perceptual distortions and misleading visualisations. All of these factors need to be considered when developing visualisation software.

Although it may not be apparent to the user, many of the techniques for visualisation inherently rely on interpolation of data (see also, Chapter 6). In other contexts, the interpolation of metrology data is undertaken with care to ensure that the algorithms used are correct and that they do not distort the data. Similar care should be used in visualisation. For example, the data in Figure 4.1 is plotted as a contour map with two different interpolation schemes. The marked points are the positions of the original data. The false nature of the faceting seen in Figure 4.1(a) is an obvious result of interpolation. However, more subtle effects can give a misleading sense of the data. Figure 4.2 shows a vector field [34] which describes the speed of movement of a rigid body moving about a fixed point. If we consider this to be the motion of a fluid and calculate stream lines using a naive algorithm, we generate the graph in Figure 4.2(b). In this graph the stream lines spiral away from the centre of rotation, which contradicts our expectation (and the correct result) that the stream lines should be circular. In this example, it is easy to predict the true form of the stream lines but that would not normally be the case with more complex data.

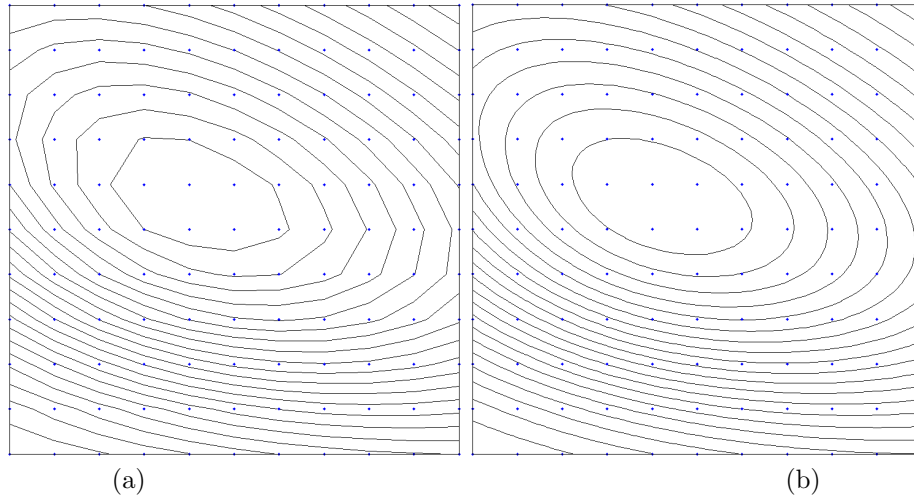


Figure 4.1: The algorithm used to interpolate data for visualisation influences the graphical form of the output: (a) linear interpolation and (b) bicubic interpolation.

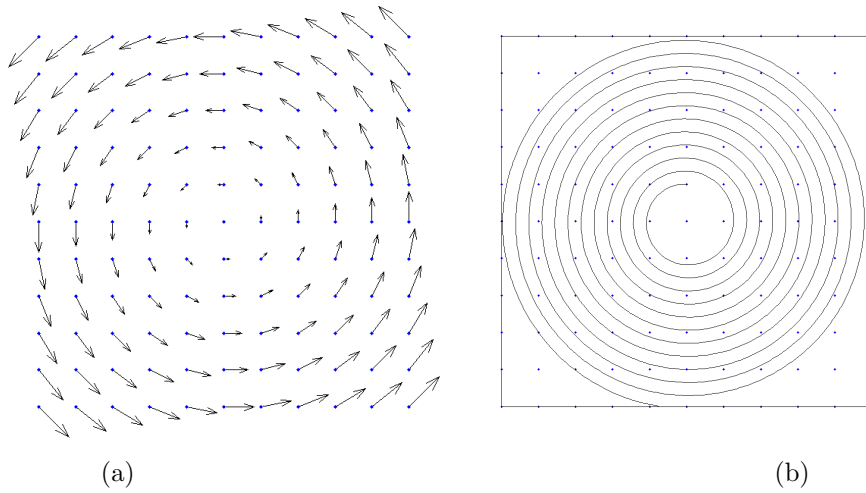


Figure 4.2: A uniformly rotating fluid has velocity vectors that are perpendicular to a line drawn from the centre of rotation and the stream lines are circular (a). The algorithm used to calculate the stream lines from this vector field gives an incorrect result (b).

4.2 Validation of visualisations

The validation of visualisation software has to address two issues:

- The manner in which the visualisation software represents and transforms the data.
- The way in which the user of the software understands and interprets the results of the visualisation.

These topics introduce challenges that go beyond those normally encountered in validating software. The canonical texts within the discipline of visualisation tend to concentrate on the second question, which is concerned with communication, i.e., how to assist the user in understanding the information being presented and to avoid misleading the viewer by creating the wrong impression of the data [28, 29, 30]. The validation process itself is not discussed. Nevertheless, it is useful to consider whether validation and testing techniques which have been developed for model validation in other areas of mathematical and software research can be adapted for use in verifying visualisations.

The first point above largely relates to the working of the visualisation software itself. The issues that need to be considered are the correctness of the algorithms used in the construction of the visualisation [21, 32, 33], including any additional mathematical operations performed on the data by the visualisation software such as interpolation or co-ordinate system transformation. The best way to establish correctness of algorithms is by carrying out formal software testing [2], so it is useful to consider whether these operations can be checked by the use of test data sets. Useful guidance on developing test data sets for validating visualisations is available [17]. Another useful check is to use a different software package or visualisation technique to look at all or part of the data for comparison.

General validation advice

The SSfM programme has produced a number of guides and reports that are of assistance when considering validation. The best-practice guide on discrete modelling [1] (incorporating an earlier guide on discrete model validation) defines the validation process as determining whether the results produced are consistent with the input data, theoretical results, reference data, etc. The key question to be considered is whether the model encapsulates all that is known about the system.

This key question has also been applied to the validation of continuous models. In particular, it is recommended that the knowledge and intuition of metrologists should be formally incorporated into the validation process [14]. The question of whether a result ‘looks right’ or ‘looks wrong’ to an experienced metrologist can be considered to be a valid test. Often, the results of such tests cannot be easily quantified. They are probably most effectively used as a simple ‘pass/fail’ criterion for the model in question, especially when the tests are essentially subjective and two metrologists may not agree, but the tests are nevertheless important. Questions that can be considered in this way concern whether the global properties of the model look right, but they can equally well be applied to visualisations of data:

- Are values the right order of magnitude?
- Do they have the correct sign?
- Are shapes, sizes and locations of key features correct?
- Is the behaviour as expected at boundaries, edges and corners?
- Do zeros and maxima appear in appropriate locations?
- Is asymptotic behaviour captured correctly?
- Do different ways of presenting the same data demonstrate consistent results?
- Are there singularities in the model and does the model behave as expected in the region of such singularities?

These questions are directly applicable to visualisations of discrete and continuous model results, and should be built into visualisation test plans. Clearly, the kinds of test described can be applied most straightforwardly to representations of geometric objects and to experimental and theoretical results that relate to geometrical objects, such as stress fields in an engineering component, or representations of electrical, acoustic, pressure or temperature fields in space. When what is being represented is a purely mathematical object, tests based on prior knowledge or intuition may be more difficult to carry out.

Perceptual problems

The computer screen, a two-dimensional surface, produces its own problems. Everyone is familiar with optical illusions which demonstrate that the human eye and brain can be deceived by images, leading to erroneous conclusions about what is being seen. Often such effects arise from tricks of perspective, in which the brain appears to use information it has about the three-dimensional world to re-interpret two-dimensional representations. An important question to consider during validation, therefore, is whether incorrect results are being masked by the way in which the observer views the image. Conversely, the results may be correct, but the observer may be being deceived by the way the brain perceives the image. The algorithms used to generate the visualisation may be correctly implemented, but a poor choice of algorithm can produce an image that misleads the observer.

Another issue that can mislead the observer is the visualisation of discrete quantities as continuous data, and vice versa [23]. The simplest example is the plotting of a smoothly varying line through discrete data points, which creates the impression that the data are continuous. Similarly, in higher dimensions, the impression almost always created by contour lines and isosurfaces is that the original data from which the contours and surfaces have been derived are continuous, even where this is not the case. In some cases, the quantities being visualised are the results of a model designed to have continuity, but if the quantities are raw measurement data then the assumption of continuity may not be valid.

The other side of this problem is the interpretation of continuous data as discrete. It can be difficult to establish whether apparent discretization in an image relates to the data being visualised or to the visualisation process itself. A simple example here is the use of contour lines or colours to represent a numerical scale (see section 3.2.2). It is important

to use sufficient colours or lines to represent the data adequately, and to ensure that specific local behaviours, which might indicate a problem with the visualisation, are not hidden by the visualisation itself. Some visualisation tools do not use smoothly-varying colours like those in Figure 3.16: instead they use a single colour to represent values between two limits. This method can encourage the observer to think of the results as being discrete disjoint sets rather than a continuum.

Another interesting example of how one might be misled by discretization has been observed by one of the authors of this report in commercial finite element software. The image in question was of an engineering component, a part of a spherical shell. If the software was asked to display the location of the nodes defined in the model or to show the shapes of the main areas of the components that made up the model, the smoothly curved nature of the surface was immediately apparent. However, when the software was asked to display the elements that made up the surface, these were represented as flat surfaces, so that the shape in question appeared to be composed of a number of facets. In fact, in the case of finite element modelling, such an observation may be helpful in that it may indicate the need for a finer mesh density for the problem of interest. This example shows the importance of a clear understanding of what is being shown: if linear elements are used to model a curved shape, then a visualisation of the elements should show facets even though the underlying shape is curved.

4.3 Recommendations

- Visualisations and the data and software used to produce them should be treated as software and normal software development good practice should be adopted.
- Algorithms used to interpolate or derive other values from the base data for visualisation should be subject to the same analysis and testing as others used within metrology.
- Use a different package or an alternative method to visualise all or part of the results.
- During visualisation, carry out informal checks that the results ‘look right’.

Chapter 5

Visualisation of Uncertainties

The main reason for tackling the issues of calculation, validation and visualisation of uncertainties is that without the consideration of these issues the effectiveness of visualisation techniques used by metrologists would be limited. Although visualisation of data or a model can give the viewer a feeling of understanding, it is impossible to know what level of assurance can be given to that understanding unless a statement or visualisation of the associated uncertainty is also available. Indeed, to the unwary, data or model visualisations using impressive graphics can be very misleading because they can tempt the viewer to ignore or forget the basic principle that every measurement has an uncertainty associated with it. This chapter considers the design and use of visualisation in the context of uncertainty evaluation, a topic of particular concern to metrologists. Further comments on the visualisation of uncertainties regarding continuous models are made in Chapter 6.

The primary document regarding uncertainty evaluation in metrology is the Guide to the Expression of Uncertainty in Measurement (GUM) [5]. It provides a framework for uncertainty evaluation based on the use of the law of propagation of uncertainty and the Central Limit Theorem. For certain problems, analytical methods may be used for uncertainty evaluation [12]. Increasingly, numerical approaches are being applied to the problem, including the Monte Carlo method [4, 10] and, for linear models of measurement, methods based on undertaking the convolution of probability distributions [20]. It is not the purpose of this Guide to describe and compare methods for uncertainty evaluation (this is done elsewhere) but rather to show how the results delivered by the methods may be visualised. In fact, the approaches to visualisation that are illustrated in this chapter are, for the most part, examples of techniques described earlier in Chapter 3 or developments of those techniques (as in the use of animation and interactive visualisation).

5.1 Objective of the visualisation

The problem of uncertainty evaluation can be divided into two phases: formulation and calculation [4, 10]. The formulation phase constitutes developing a model of the measurement and using information about the input quantities to the model to assign probability distributions to the values of those quantities. The purpose of the calculation

phase is then to obtain a probability distribution for the value of the output quantity from the model, and from it to form an estimate of the output quantity value and the associated uncertainty and a coverage interval for the value. In the context of uncertainty evaluation, visualisation may be used to support both phases. In particular, the objectives of visualisation include:

- Communicating information about the input quantities to the measurement model as part of the formulation phase of uncertainty evaluation.
- Communicating information about the output quantity obtained from the calculation phase of uncertainty evaluation.
- Comparing the results obtained from different approaches to the calculation phase of uncertainty evaluation (e.g., analytical, the law of propagation of uncertainty and the Monte Carlo method).
- Comparing (different) information about the same quantity obtained from different sources, e.g., where a quantity is measured using different instruments or by different organisations/laboratories as may arise in interlaboratory comparisons.

The ‘objects’ to be visualised include:

- The estimate of a quantity with the associated standard uncertainty.
- The estimate of a quantity with a coverage interval for the value of the quantity corresponding to a stated coverage probability.
- The probability distribution for the value of a quantity described by a probability density function (PDF) or the corresponding distribution function (DF) or by (finitely-many) samples drawn from the PDF for the value of the quantity.¹
- All the above, but for vector-valued quantities.
- All the above, but regarded as ‘functions’ of ‘parameters’ of the problem. The ‘parameters’ may include:
 - Frequency or wavelength, as arises for spectrally-varying quantities.
 - The estimates of the values of the input quantities.
 - The standard uncertainties (and mutual standard uncertainties) associated with the estimates.
 - The measurement model, e.g., unweighted mean, weighted mean and median as may arise in the analysis of data from key comparisons [8].

5.2 Designing the visualisation

In this section a number of visualisation techniques that are useful in the context of uncertainty evaluation are presented. The visualisation techniques are illustrated using a

¹The representation of the probability distribution by finitely-many samples arises, for example, from the application of the Monte Carlo method to the calculation phase of uncertainty evaluation.

particular problem in uncertainty evaluation concerning the transformation from Cartesian to polar representations of a complex-valued quantity. The formulation and calculation phases of this problem are first described.

5.2.1 Co-ordinate system transformation

An issue that arises in some areas of metrology (e.g., electrical, acoustical, etc.) is that measurements are typically made of complex-valued quantities. For some purposes these quantities are represented in terms of their real and imaginary parts (Cartesian co-ordinates) and for other purposes in terms of magnitude and phase (polar co-ordinates). Uncertainties are invariably associated with the measurements. In particular, measurement results frequently have to be transformed from one system of co-ordinates to the other. The associated uncertainties have accordingly to be transformed in a valid manner.

The model of measurement relating input quantities $(X, Y)^T$ (Cartesian co-ordinates) to output quantities $(R, \Theta)^T$ (polar co-ordinates) is

$$R^2 = X^2 + Y^2, \quad \tan \Theta = Y/X. \quad (5.1)$$

This is an *implicit* model for R and Θ [10]. Since R is a magnitude, and is known to be non-negative, it is determined uniquely by considering the positive square root of R^2 . By using a ‘four quadrant inverse tangent function’ (often denoted by ‘atan2’), Θ is determined uniquely from the model to satisfy $-180^\circ < \Theta \leq 180^\circ$.

Estimates $(\mu_x, \mu_y)^T$ of the values of Cartesian co-ordinates $(X, Y)^T$ are available. The uncertainty matrix (covariance matrix)

$$V_{x,y} = \begin{bmatrix} \sigma^2 & 0 \\ 0 & \sigma^2 \end{bmatrix}, \quad (5.2)$$

where σ denotes the standard uncertainty associated with μ_x and μ_y , and the estimates are uncorrelated, is assumed to be available. A bivariate Gaussian distribution with expectation $(\mu_x, \mu_y)^T$ and uncertainty matrix $V_{x,y}$ is assigned to the values of $(X, Y)^T$, e.g., on the basis of maximum entropy considerations [10, 39, 40].

The bivariate distribution for the values of the polar co-ordinates $(R, \Theta)^T$ is required. From this distribution the expectation $(\mu_r, \mu_\theta)^T$ of the values of $(R, \Theta)^T$, the uncertainty matrix $V_{r,\theta}$ associated with $(\mu_r, \mu_\theta)^T$, and a coverage region for the values of $(R, \Theta)^T$, corresponding to a stipulated coverage probability, can then be determined.

A treatment of the calculation phase corresponding to the above formulated problem is available [9]. The treatment covers (a) the application of the law of propagation of uncertainty, (b) the application of the Monte Carlo method, and (c) the derivation of an analytical solution to the formulated problem. The analytical solution and the application of the Monte Carlo method are summarised below, and used as the basis of the visualisations subsequently described.

Analytical treatment

An analytical treatment [9, Appendix D] of the co-ordinate system transformation problem formulated above gives

$$g(R, \Theta) = \frac{R}{2\pi\sigma^2} \exp\left(\frac{-(R \cos \Theta - \mu_x)^2 - (R \sin \Theta - \mu_y)^2}{2\sigma^2}\right) \quad (5.3)$$

for the joint PDF for the values of R and Θ , from which the marginal distributions for the values of R and Θ may also be derived.

It is shown [9, Appendix D] that when $\mu_x = \mu_y = 0$, the output quantities are mutually independent, with the value of R described by a χ -distribution with two degrees of freedom,² or equivalently a Rayleigh distribution [24], with PDF

$$g(R) = \frac{R}{\sigma^2} \exp\left(\frac{-R^2}{2\sigma^2}\right), \quad R \geq 0, \quad (5.4)$$

and the value of Θ by a rectangular distribution with PDF

$$g(\Theta) = \frac{1}{360}, \quad -180^\circ < \Theta \leq 180^\circ. \quad (5.5)$$

Monte Carlo method

A key task in the application of the Monte Carlo method to the problem formulated above is to be able to draw random samples $(x_k, y_k)^T$, $k = 1, \dots, M$, where M is the number of Monte Carlo trials, from a bivariate Gaussian distribution with expectation $(\mu_x, \mu_y)^T$ and uncertainty $V_{x,y}$ given by expression (5.2) [4, 11]. For each random sample $(x_k, y_k)^T$ so obtained, the model (5.1) is used to determine corresponding values $(r_k, \theta_k)^T$ of the output quantities, viz.,

$$r_k^2 = x_k^2 + y_k^2, \quad \tan \theta_k = y_k/x_k, \quad k = 1, \dots, M.$$

5.2.2 Slices

Figure 5.1 shows the PDFs corresponding to the marginal distributions obtained from an analytical treatment of the co-ordinate system transformation problem for the values of magnitude R and phase angle Θ for particular values of μ_x , μ_y and σ .

In this example, the ‘U’-shape of the PDF for the value of Θ is a consequence of the choice of $\pm 180^\circ$ as the point at which we choose to ‘wrap’ the phase angle. The choice of 0° for this point would provide a visualisation of the PDF that would show the peak of the PDF more clearly. The display of marginal PDFs as in Figure 5.1 necessarily limits the information about R and Θ that is communicated: the graphs do not show how the values of R and Θ vary as functions of the parameters of μ_x , μ_y and σ , or the extent to which the quantities R and Θ are mutually dependent. More complex visualisations, considered below, are necessary for this purpose. It is common to add to such graphs the position of

²The square root of the sum of the squares of two standard Gaussian variates.

the expectation of the quantity represented by the PDF as well as the endpoints of a coverage interval for the value of the quantity. Although not communicating as much information about the quantity as a PDF, it is in terms of such summary statistics that the results of an uncertainty evaluation are typically reported.

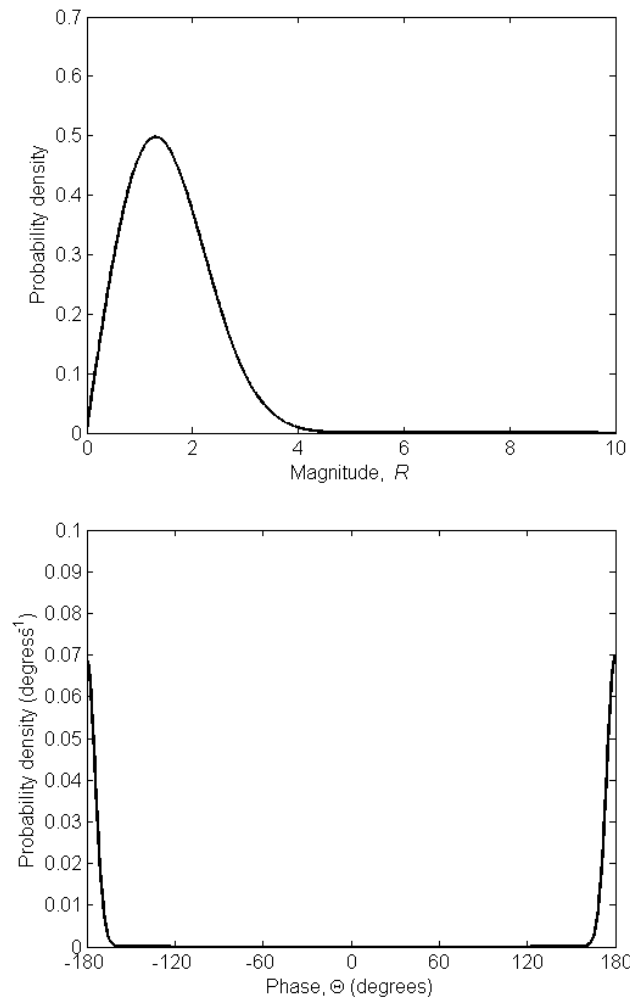


Figure 5.1: The marginal PDF for the value of magnitude R corresponding to $\mu_x = -1$, $\mu_y = 0$ and $\sigma = 1$, and (bottom) that for the value of phase angle Θ , corresponding to $\mu_x = -10$, $\mu_y = 0$ and $\sigma = 1$.

5.2.3 Elevated surfaces

The stacking (vertically or, here, side-by-side) of a sequence of slices provides a representation of a marginal PDF in the form of a surface for which the ‘independent variables’ are the output quantity (here R or Θ) and the ‘slice’ variable (here μ_x). It provides a visualisation of the marginal PDFs for the values of R and Θ using a two-dimensional display, either as contour plots or as elevated surfaces. Figure 5.2 shows

these marginal PDFs as *elevated surfaces*. A two-dimensional display such as an elevated display depicts the overall shape of the ‘complete’ function. The height of the surface at any point is proportional to the value of the function depicted by the surface at that point. There are likely to be ‘hidden regions’ (which can sometimes be revealed using a different viewing point). The display may or may not incorporate perspective effects. The consequent influence on the perceived ‘size’ of surface features should be taken into account.

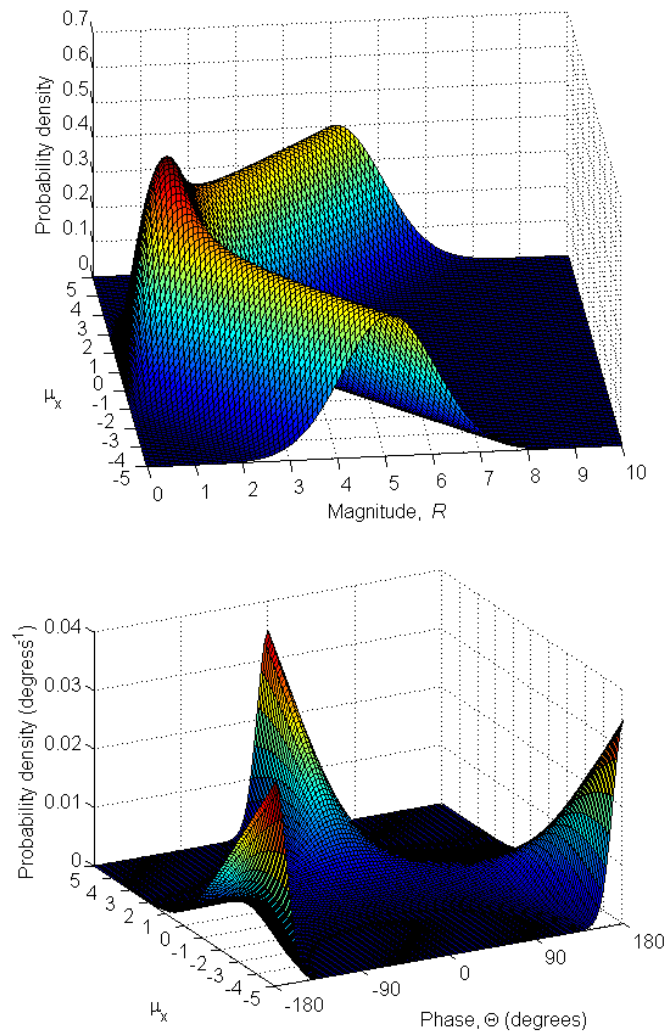


Figure 5.2: Elevated surfaces showing, as a function of μ_x , the marginal PDFs for the value of R and (bottom) those for the value of Θ .

5.2.4 Movies: slices, elevated surfaces and contour plots

The presentation of a succession of slices (in time) as one of the defining parameters changes throughout its permissible range of values constitutes a movie. It provides an alternative to the use of elevated surfaces to portray the marginal PDFs for the values of R and Θ as μ_x cycles through a sequence of values.

A movie of the marginal PDFs for the value of R with $\mu_y = 0$ and $\sigma = 1$, as μ_x ranges from -10 to 10 in steps of 0.5 , can be viewed from the electronic form of this document (see Appendix B). A movie of the marginal PDFs for the value of Θ for the same values of μ_y and σ , and the same range of values of μ_x , can also be viewed. It is emphasized that each frame of these movies shows the marginal PDF for the value of magnitude or phase angle corresponding to a particular value of μ_x (as in Figure 5.1).

Movies may also be used to show the *evolution* of elevated surfaces. A movie of the joint PDFs for the values of $(X, Y)^T$ and $(R, \Theta)^T$ with $\mu_y = 0$ and $\sigma = 1$, as μ_x ranges from -5 to 5 in steps of 0.5 , can be viewed. On the left of the movie is shown the evolution of the joint PDF for the value of $(X, Y)^T$ (a bivariate Gaussian distribution) and on the right that of the corresponding PDF for the value of $(R, \Theta)^T$ (given by expression (5.3)). A movie of the joint PDFs represented as contour plots, in place of elevated surfaces, can also be viewed.

A movie (in the context considered here) constitutes a sequence of adjacent slices through a surface. Each frame of the movie (corresponding to a slice) displays the actual shape of that slice. To obtain a complete understanding of the surface, however, requires the viewer to ‘remember’ the whole sequence. Unless the variable being animated is time, the association of frames with a changing parameter rather than time can be difficult to appreciate. A movie emphasizes the variation with a particular choice of variable (μ_x , e.g., as here). A more complete understanding would be obtained also by an additional movie involving ‘the other variable’. Limited interaction is possible through changing a parameter on which the movie depends, and re-viewing.

5.2.5 Scatter plots

Some surfaces, especially approximations to PDFs generated using the Monte Carlo method, are represented by point clouds (unless further processing is used to provide, usually, continuous approximations). Figure 5.3 shows 10 000 points sampled randomly from a PDF given by the product of two Gaussian PDFs, for the values of X and Y , each having mean zero and unit standard deviation, and the corresponding points in the magnitude-phase plane. The graphs shown in Figure 5.3 provide a visualisation of the raw (un-processed) output obtained from the application of the Monte Carlo method.

5.2.6 Scaled histograms

For the case that provides the data shown in Figure 5.3, the joint PDF for the value of R and Θ can be expressed as the product of the PDF for the value of R and that for Θ . Thus, an approximation to the PDF for the value of R can be provided by presenting the (10 000) R -values obtained as a ‘scaled histogram’, with a similar statement for Θ .

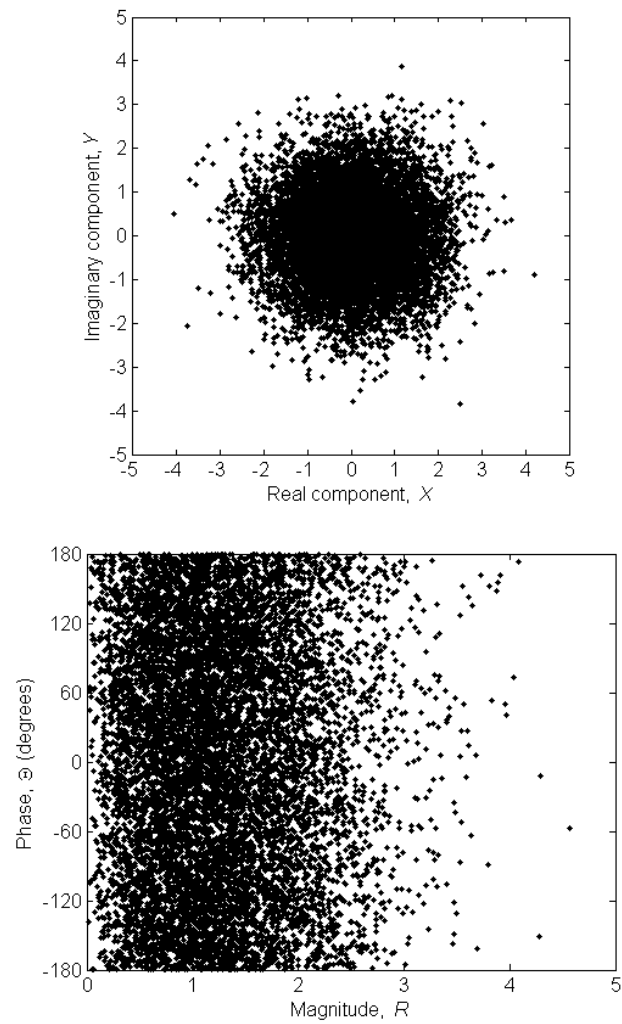


Figure 5.3: 10 000 points sampled randomly from a PDF given by the product of two Gaussian PDFs, for the values of X and Y , each having mean zero and unit standard deviation, and (bottom) the corresponding points in the magnitude-phase plane.

Figure 5.4 shows such an approximation to the PDF for the value of R . The range of R -values is divided into 20 sub-intervals of equal width. The area of a rectangle of the histogram is proportional to the number of R -values within the corresponding sub-interval, with the resulting histogram scaled to have an area of unity. The figure also shows the exact PDF, a Rayleigh distribution as given by expression (5.4). Figure 5.4 also shows the counterpart for Θ , for which the exact PDF is a rectangular distribution as given by expression (5.5).

The graphs shown in Figure 5.4 display information about R and Θ obtained by processing the data presented in the bottom graph of Figure 5.3. It is important to make a sensible choice for the number of sub-intervals used to define the scaled histograms, otherwise important information about the quantities may be lost. An alternative is to display the output obtained from an application of the Monte Carlo method as a distribution function [4]. The advantages of this are that it is not necessary to make a choice of the number of sub-intervals, and the resulting graph is often ‘smoother’ than the corresponding scaled histogram. A disadvantage is that it can be more difficult to extract information, such as skewness of the distribution, from the distribution function, as well as to compare (visually) different distribution functions.

5.2.7 Interactive visualisation

The above visualisations are not under the direct control of the human observer. Of course, a mild form of interaction is possible through the provision by the user of particular values for μ_x , e.g., and the examination of the corresponding display. An *interactive* visualisation has the advantage that the observer has the opportunity to display the ‘object’ in a manner that is deemed appropriate to the application. The use of a limited set of viewpoints may paint an unfaithful picture of the surface, however. Parts of the object are likely to be obscured by other parts. For both reasons a detailed exploration using the provided facilities (roll, turn, zoom, etc.) is advised. For purposes of publication of results, one or more of the above forms of display may be more suitable. If, however, the publication is in the form of an electronic document such as the present document when viewed online, interactive visualisation through that document is possible.

For example, three-dimensional models of the elevated surfaces shown in Figure 5.2 are provided in the electronic form of this document (see Appendix B) for the marginal PDFs for the value of R and those for the value of Θ . In these models, the axes are labelled with ‘mux’ to denote μ_x , ‘r’ the value of the magnitude R , ‘theta’ the value of the phase angle Θ , and ‘PDF’ to denote probability density (dimensionless for magnitude R and in units of ‘degrees⁻¹’ for phase angle Θ). These models can be examined from different viewpoints using the controls on the browser in accordance with the type of movement required by the user. Doing so constitutes a simple form of interaction but, unlike the interactive visualisation described below, there is no provision for the values of the parameters μ_y and σ that determine the object displayed to be altered.

An interactive tool, having several user controls, for visualising co-ordinate system transformations is also provided in the electronic form of this document. The interactive tool is described below.

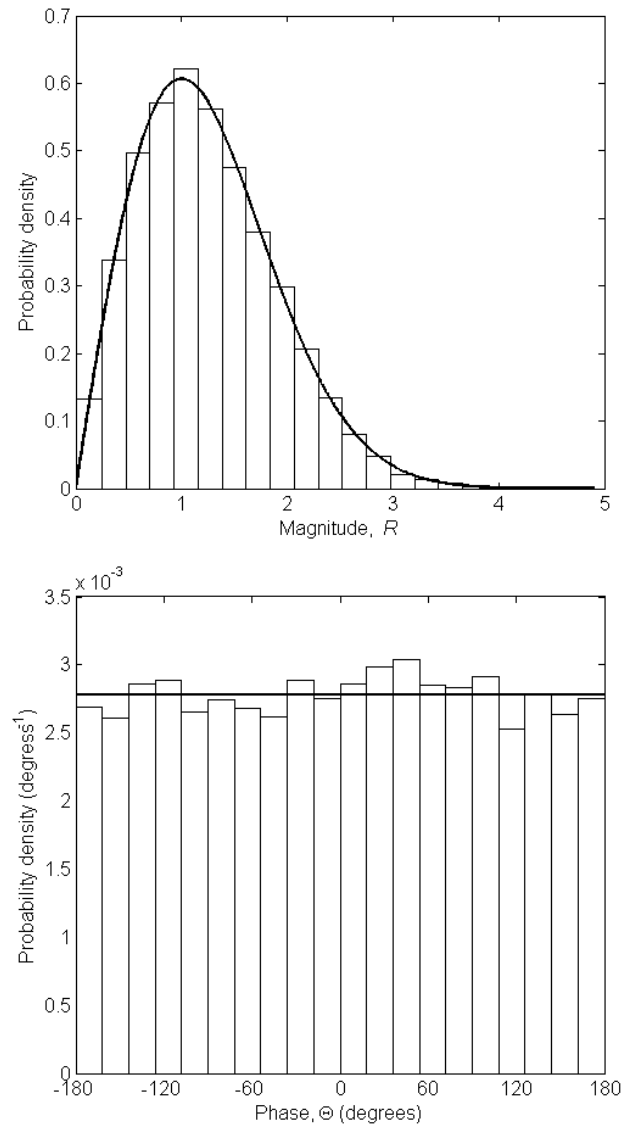


Figure 5.4: An approximation in the form of a scaled histogram to the PDF for R and the analytic PDF for R , a Rayleigh distribution, and (bottom) the corresponding approximation to the PDF for Θ , with the exact PDF for Θ , a rectangular distribution.

Operating the visualisation

Figure 5.5 contains a screen shot from the tool. The screen is subdivided into two parts. The left side of the tool relates to the model input quantities X and Y . It depicts a circular coverage region, shown as a red disc, corresponding to a required coverage probability p , for the values of X and Y . It also shows the controls that are used to choose μ_x , μ_y and σ . μ_x and μ_y can be selected using the sliders or entered directly into the appropriate control boxes. σ is also entered using a control box on the lower left side of the tool. Of the two remaining control boxes on the lower left of the tool, 'coverage probability (%)' is used for p , and *Grid size* for the resolution to which the transformation is to be carried out. Higher values of *Grid size* result in more accurate results, but in longer computation times.

The right side of the tool relates to the output quantities R and Θ . It shows as a three-dimensional object a $100p$ % coverage region for the values of R and Θ displayed in a cylindrical co-ordinate system. This object can be rotated, transformed, etc., as required, using the standard controls on the right part of the screen. If the cylinder is clicked, it will unroll and show a flat representation of the cylindrical co-ordinate system with R in the vertical and Θ in the horizontal direction. Clicking on this plane will cause it to roll up into a cylinder again.

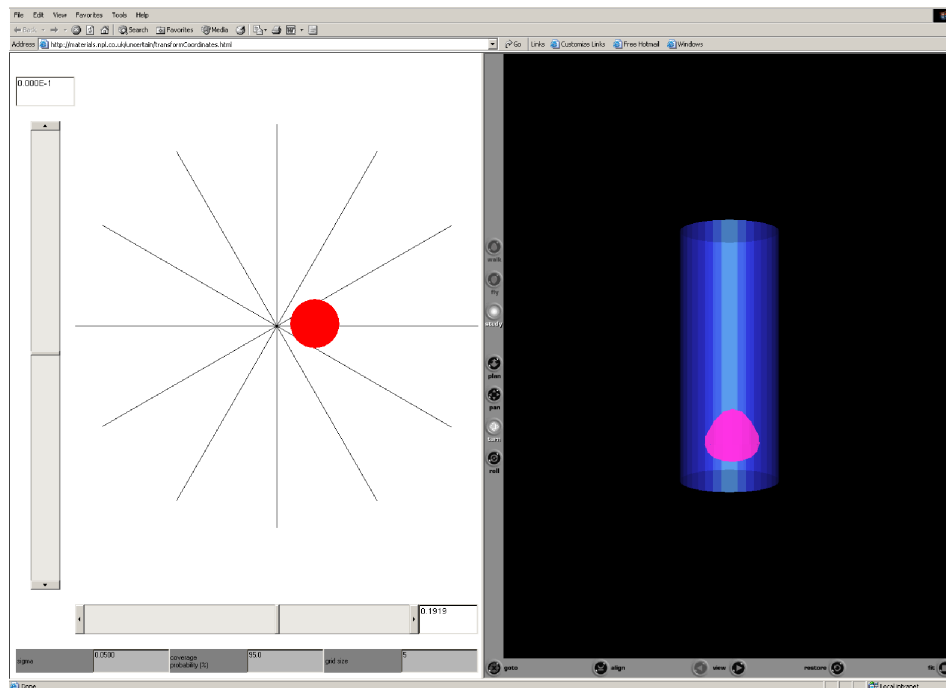


Figure 5.5: Screen shot of the tool for visualising co-ordinate system transformations.

5.2.8 Parallel co-ordinates

The parallel co-ordinate technique is useful for exploring the relationships between the summary statistics of the different quantities. The data set used for this exploration was

created by choosing 41 values of μ_x between -10 and 10, 41 values of μ_y between -10 and 10, and 50 values of σ between 0.1 and 5, and then running a Monte Carlo simulation for every possible combination of these values, as described in section 5.2.1 (84050 simulations in total). Each simulation is run for $M = 1000$ trials, and the values of R and Θ obtained are summarised by their sample means m_r and m_θ , sample standard deviations s_r and s_θ , and the sample correlation $s_{r,\theta}$.

The resulting data set has eight dimensions ($\mu_x, \mu_y, \sigma, m_r, s_r, m_\theta, s_\theta, s_{(r,\theta)}$): note that distribution values rather than sample values have been used for μ_x, μ_y , and σ), and 84050 points. The data have been visualised using Curvaceous Software's CVE, a parallel coordinates visualisation software tool. The tool enables data exploration by allowing the user to highlight subsets of the data in a different colour. Users can choose the highlighted subsets in several different ways. The plots used in this section have all been created by limiting the range of one or more of the dimensions, with the limit ranges on the plots indicated by red triangles. In all plots the axes, from left to right, represent the quantities $\mu_x, \mu_y, \sigma, m_r, s_r, m_\theta, s_\theta, s_{(r,\theta)}$.

A number of interesting points have been noted from the visualisations. Figure 5.6 shows that for a fixed range of m_r , the line segments linking the μ_x and μ_y axes have the characteristic shape associated with a circular region (further investigation shows that the lines represent an annular region, as would be expected). Figure 5.7 shows a similar visualisation highlighting a small range of s_r . The highlighted line segments show that a small value of s_r requires a small value of σ , independent of μ_x and μ_y . Further investigation shows that s_r is proportional to σ , and for a fixed value of σ , s_r is approximately proportional to m_r .

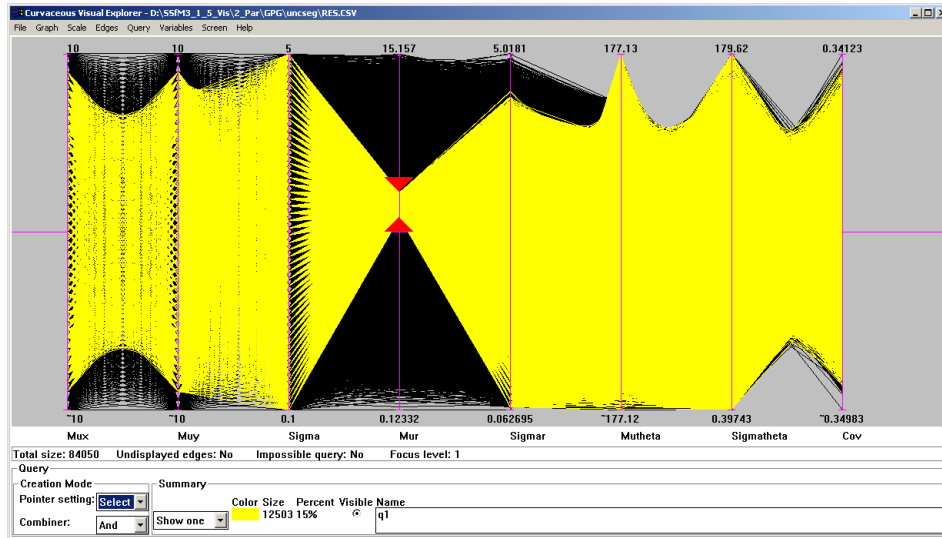


Figure 5.6: Highlighted line segments represent a fixed range of m_r .

Figure 5.8 highlights values with an abnormally large value of s_θ (the maximum value is approximately 180°). The corresponding values of μ_x are all negative, and the values of μ_y are all small. Hence these data points correspond to the situation where the (x_i, y_i) are grouped around the negative half of the x axis. The definition of Θ given in equation 5.1 means that the value of Θ has a jump discontinuity from -180° to 180° as it crosses the

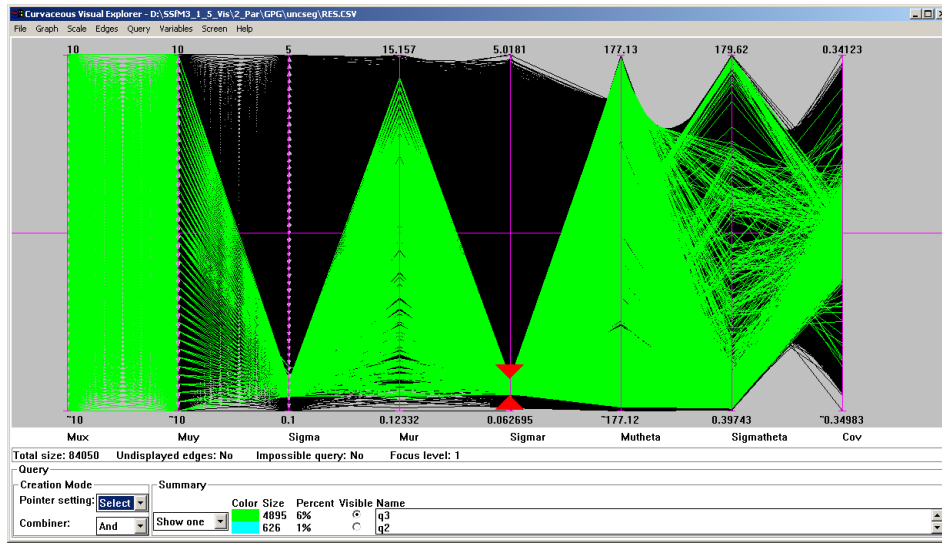


Figure 5.7: Highlighted line segments represent a fixed range of s_r .

negative half of the x axis. Hence any Monte Carlo simulation in which the (x_i, y_i) are centred around the negative half of the x axis will include some large negative values of Θ and an approximately equal number of large positive values, with the result that the mean value m_θ will be smaller in absolute value than any of the values of Θ used to calculate it. Therefore the value of s_θ will be abnormally large, due to the absolute values of the θ_i , and the small size of the mean value m_θ .

This example illustrates one of the potential pitfalls of summarising distributions using the mean and variance, particularly if the statistics were obtained using Monte Carlo simulation. A typical full distribution is shown in figure 5.1, showing that if the distribution is wrapped round so that -180° and 180° meet, the variance s_θ is quite small. The use of Monte Carlo simulation in a region where one of the result quantities is discontinuous is not valid: if the technique is to be used for negative values of μ_x and small values of μ_y , then Θ needs to be redefined as lying between 0° and 360° , for example.

Visualisations were also created to investigate the correlation coefficient $s_{(r,\theta)}$ but the coefficient values suggest that there are no strong correlations between the data, so no plots have been shown.

5.3 Recommendations

- The probability distribution for a quantity conveys more information about the quantity than summary statistics obtained from the distribution. However, it is often summary statistics (expectation, standard deviation and coverage interval) that are reported as the result of an uncertainty evaluation.
- The joint probability distribution for a vector-valued quantity conveys more information about the quantity than the marginal probability distributions for the individual components.

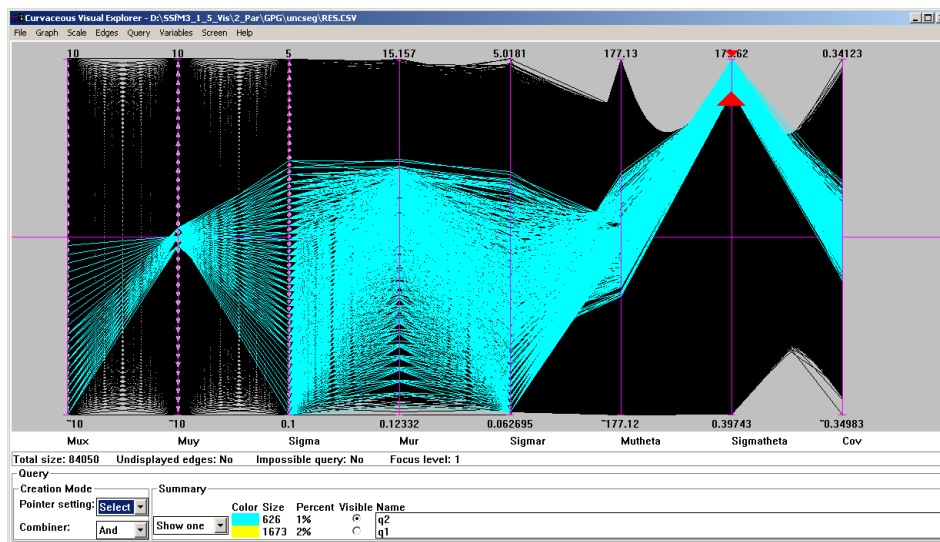


Figure 5.8: Highlighted line segments represent a fixed range of s_θ .

- Choose a visualisation technique that is suitable for the available information: this may take the form a PDF or DF or randomly-drawn samples from a PDF.
- Consider using elevated surfaces, movies and interactive visualisations to allow a user to explore the dependence of the result of an uncertainty evaluation on parameters of the problem.

Chapter 6

Visualisation in Continuous Modelling

Much continuous modelling software, such as that used for finite element and finite difference analysis and computational fluid dynamics, incorporates visualisation for displaying both a model and the results obtained from the model. This chapter of the Guide, summarising previous work [41], provides advice on the use of visualisation in continuous modelling.

6.1 Continuous modelling results

For many users of continuous modelling software it is likely that their first attempts to understand the results obtained from the software will be through a visualisation of those results rather than of the numerical values themselves. In the context of metrology, much continuous modelling is undertaken to assist in understanding physical processes, to help design experiments, and to identify and quantify systematic effects in experiments. Thus, the model and its visualisation may not be an end in themselves, but rather one aspect of a much larger investigation.

Visualisation will appear to many software users to offer a succinct and cost-effective way of studying the results obtained from a model. However, it must be recognised that the use of visualisation to understand numerical results incorporates an additional level of processing between the results obtained from the model and the user. It may not always be obvious to the user, especially in the case where ‘black box’ software packages are used, what are the assumptions and limitations associated with the visualisation.

Visualisation in continuous modelling is also commonly used for the investigation of general questions about model results [41], such as:

- Do the results ‘look right’?
- Where are the maxima and minima?
- How do the results from different models compare?
- Over what range of values do the results vary?
- What are the broad trends in the results?

Ideally, the visualisation of results should not add any further inaccuracy to that inherent in the modelling process. At the visualisation stage, the model results are available and should not be altered by the visualisation. If the model results can be displayed sufficiently accurately and clearly, the viewer will gain a good understanding of them and be in a position to draw reliable conclusions. However, in reality this does not always happen, as other authors have shown [32].

6.2 Interpolation, extrapolation, and smoothing

Many algorithms for solving continuous modelling problems consider a discretised version of the problem. The algorithms deliver results at a number of points (known as mesh points) within the region of interest. Often results are also calculated at locations other than the mesh points. It is particularly important to note the discrete nature of the results (corresponding to a discrete set of points) because common methods for visualising the results, such as using contour plots, give the impression of the results as continuous quantities. Additionally, many such visualisations use smoothing or local averaging to produce a more aesthetically pleasing display. Smoothing ensures continuity of results across element boundaries, but can therefore hide discontinuities and erroneous values.

The concerns raised here are similar to those that arise when creating a line plot from discrete data. A scatter plot of discrete points imposes no assumptions about the behaviour of the data, but can be difficult to interpret. Joining neighbouring points with straight lines makes understanding easier but may be misleading. Similarly, fitting a higher-order polynomial curve can give a better interpretation of behaviour but imposes strong assumptions about the data.

Some particular issues regarding interpolation, extrapolation and smoothing in the context of continuous modelling are discussed below.

6.2.1 Interpolation and extrapolation

When a finite element model is developed, some functional approximation is made to describe the local behaviour of the solution. It is reasonable to assume that the solution can be visualised using the same functional approximation, and derivative quantities such

as gradients, stresses, and strains can be visualised using the derivatives of the approximation.

For example, if a stress analysis model is developed using linear elements, then the displacements should be visualised as linear, and the stresses and strains as constant, within each element. This choice of visualisation will lead to discontinuities in the visualisation of stresses and strains, but it will mean that the visualisation is consistent with the model results and does not impose additional assumptions. Most proprietary finite element packages carry out this type of visualisation automatically.

Other techniques, such as finite difference methods and finite volume methods, do not involve interpolation in their derivation, and so more careful consideration needs to be given to the display of the results obtained from these methods.

Finite difference methods use Taylor expansions to approximate derivatives of the variable of interest and hence produce a system of equations involving the values of the variable at the mesh points. These methods assume local continuity of the variable and some of its derivatives, but do not impose other restrictions on the values between the mesh points.

Finite volume methods rely on the assumption of constant variable values within a small volume surrounding each mesh point, but the methods are frequently further complicated by the use of different meshes for different variables. For instance, the application of a finite volume method to solving Maxwell's equations often uses one mesh for the magnetic field and a different, offset, mesh for the electric field. The use of two meshes in this way means that careful interpretation is needed, particularly if a derived quantity that is a function of both magnetic and electric field is of interest.

The lack of assumptions about variable behaviour made by these methods means that an appropriate way of visualising the results is to regard each mesh point as being representative of an element and to assume that the result is constant within that element as shown in Figure 3.23(b). An example to show the application of this technique to visualise the results obtained from a finite difference model is available [41]. The main difficulty with assuming quantities are constant within an element is knowing how to assign (and display) a value for the quantity at points at the interface between two elements. There are two cases to consider, corresponding to continuous and discontinuous quantities. If the quantity is expected to be continuous and the model is adequate, then the values for the two elements should be sufficiently close that an average of the values can be used. Moreover, the differences between the element values and (hence) their average should be small. If the quantity is expected to be discontinuous at the interface, then both values should be presented together with an indication of which result lies on which side of the interface.

6.2.2 Smoothing

Smoothing is usually applied to quantities that have been evaluated at internal points, and involves extrapolating these quantities to the mesh points using an appropriate approximation, and then taking a local average to produce a single value at each mesh point and local continuity.

The main problem with smoothing is that it can disguise discontinuities, which can lead to

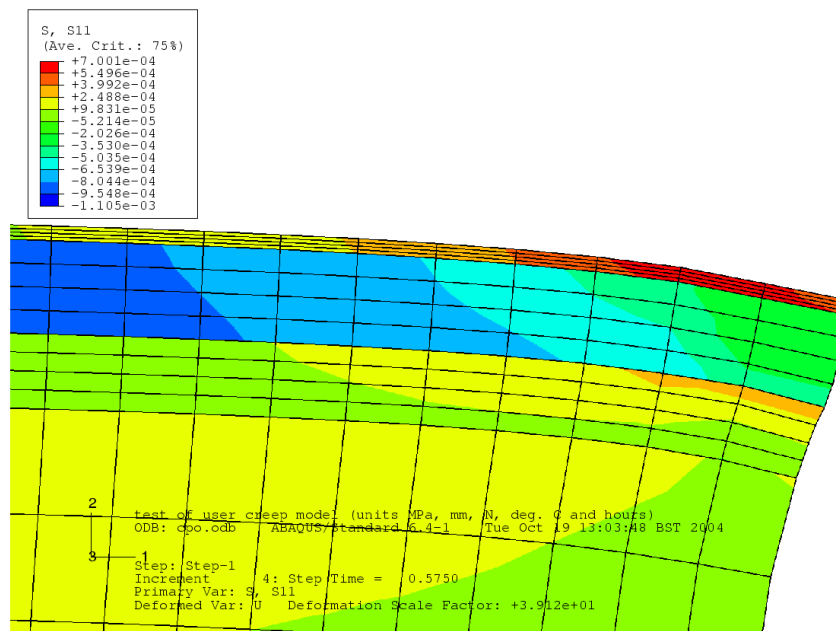
errors being hidden. Sometimes a quantity, such as stress, is expected to be continuous but has not been constrained to be continuous in the model. A visual check of the smoothness of the results for such a quantity can be used as a model validation method [14]. If smoothing is applied when visualising such a quantity, the results may appear to be continuous when they are not.

There are some models where discontinuities are expected to occur. For instance, in a stress analysis of a structure containing two materials with different elastic moduli, it is expected that the normal stresses will be discontinuous across any join between the two materials. If the displayed results are to show these discontinuities, smoothing must be deactivated. This is shown in Figure 6.1, which shows two plots of the normal stress created by loading a sample made up of several different materials. The unsmoothed plot, (a), includes features that the smoothed plot, (b), does not.

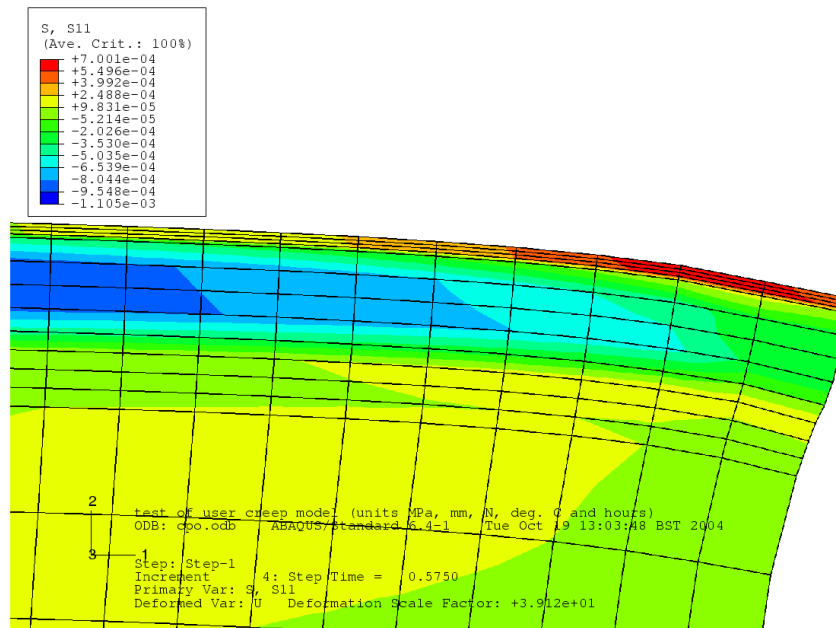
For the majority of well-behaved models that do not contain discontinuities in material properties or boundary conditions, the smoothed and unsmoothed results should look broadly the same since in many cases physical quantities and their gradients are expected to be continuous. A visual inspection of the differences between smoothed and unsmoothed results can, therefore, provide insight into the validity of the model results.

6.3 Visualisation parameter choices

The increasing provision of proprietary finite element packages for people with little or no mathematical background has led to improvements in the ease of use of the post-processing and visualisation parts of such packages. Recent increases in the range of problems that finite element packages can solve has led to improvements in the range of options offered by the corresponding visualisation software. These trends have meant that many packages now have a large number of visualisation parameters, many of which are set to default values of which the user may not be aware. Although the default values are generally chosen to be suitable for the majority of models, they should be checked to ensure that the visualisation is displaying what the user expects.



(a)



(b)

Figure 6.1: Smoothing can obscure features of the results. An unsmoothed plot of normal stress, (a), has features that the smoothed version, (b), does not.

The following examples of visualisation parameters will be discussed further below:

- Choice of viewing surface and clipping planes.
- Choice of local or global co-ordinate systems for vector and tensor quantities.
- Scaling factors.
- Display tolerances.

In general, the documentation for the visualisation software will include descriptions of all the parameters and their default values.

6.3.1 Choice of viewing surface

Some finite elements may have several integration points through their thinnest dimension despite that dimension being theoretically negligible. This occurs, in for example beams with non-uniform cross sections, shell elements, and elements used for modelling laminates. These points are necessary to describe the through-thickness variation of the results, but since the elements are visualised as one- or two-dimensional objects, the through-thickness variation is not displayed. Instead, the user chooses the results from a single set of integration points for display. The integration points are usually identified with the upper, middle or lower surface of the structure, with the default option often being the middle surface.

Similarly, if a user wishes to examine results from mesh points in the interior of a three-dimensional object, a clipping plane must be defined. A clipping plane defines a cross-sectional slice through the structure so that the interior of the model can be seen. The software does not display any parts of the model that lie in front of the clipping plane, in order that the required results can be seen.

Both these features enable users to view otherwise inaccessible results, but the default values controlling them can be misleading. For example, consider the plots shown in Figure 6.2. These plots show the stress generated in the first vibrational mode of a plate clamped at one end. Figure 6.2(a) shows the default display: the stress is virtually zero throughout the plate, which is not what the user would expect. These results are from the central plane of the plate, which is a neutral surface of the problem. Figure 6.2(b) shows the stresses on the top surface of the plate, which are of the order of magnitude that the user expects.

6.3.2 Choice of local or global co-ordinate systems

Many quantities calculated using continuous models are vector or tensor quantities, for example, stresses, strains, electro-magnetic fields and velocities. Such quantities are often displayed as contour plots by resolving them into component form relative to some co-ordinate system.

Many types of finite element are defined in terms of local co-ordinate systems. For example, it is common practice to take the direction lying along the length of a beam to be

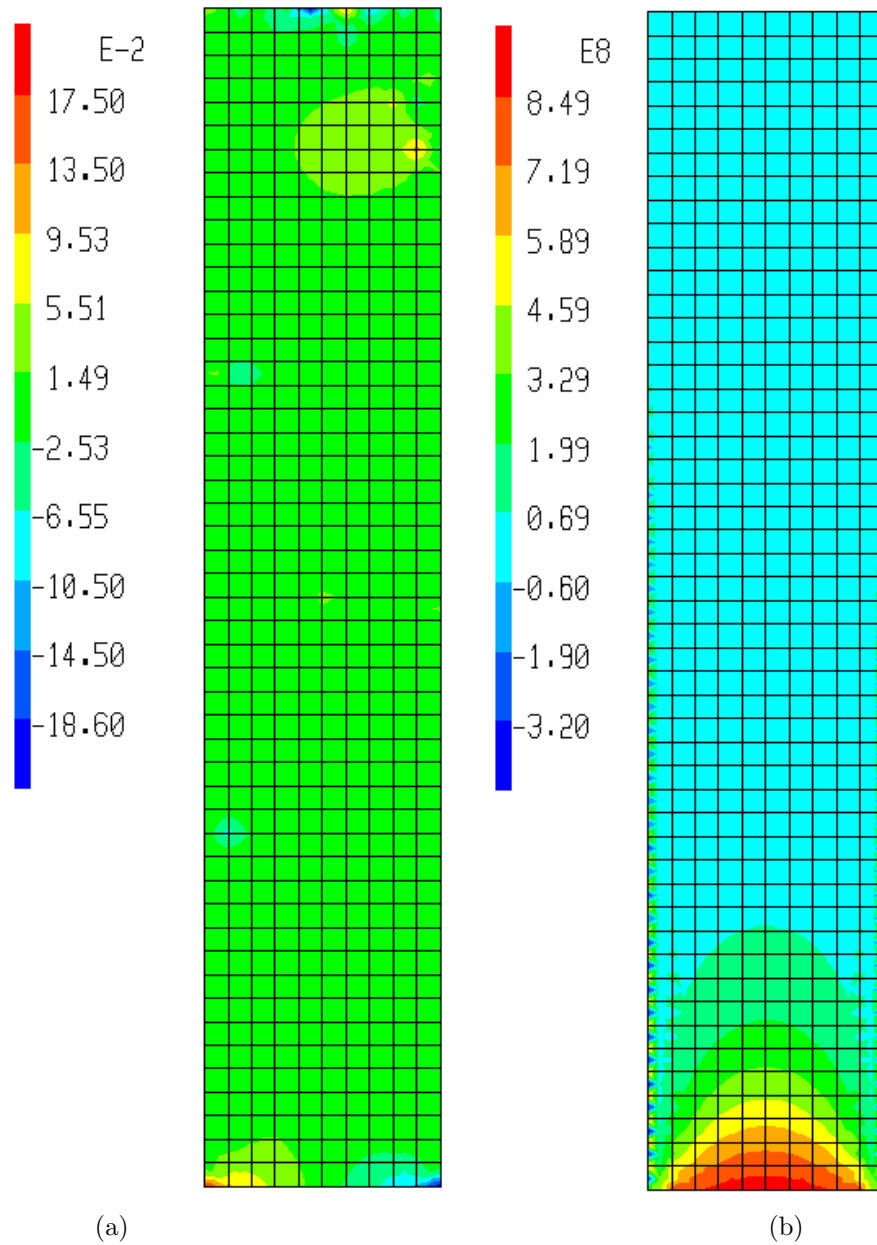


Figure 6.2: Stress contour plots from (a) the middle and (b) the top surfaces of a vibrating plate.

its local x direction, no matter which way it is oriented relative to the global co-ordinate system. Similarly, when defining complicated geometries it is often convenient to define local co-ordinate systems, such as local curvilinear axes, in order to construct the geometry more accurately.

The existence of different co-ordinate systems means that there can easily be confusion as to which co-ordinate system has been chosen for the display of vector and tensor quantities. In some cases the confusion can be removed by plotting vectors as arrow plots, or by plotting tensors as principal values on a set of principal axes. However, not every package offers these options and such plots do not always clarify the situation, particularly in cases where principal axis directions change by large amounts between adjacent elements. Hence it is important to understand which axes are being used when visualising vector and tensor quantities.

6.3.3 Scaling factors

Visualisations of the results of continuous models are affected by scaling factors in the same way as those of discrete models (see section 3.2.7). Many continuous models simulate large structures that are displaced by a small amount. For instance, a probe tip may be deflected by a few micrometres, but the probe itself may be hundreds of thousands of micrometres in size. This difference in scale means that if the deflection is shown to the same scale as the model geometry, there will be no difference between the visualisation of the undeflected and deflected structures. If one particular part of the model is of interest, visualising only a small detail of the model (perhaps by ‘zooming into’ the visualisation) provides one approach to avoiding the misleading effect on the visualisation of different scales.

6.3.4 Display tolerances

Some continuous models produce more results than a visualisation package can process. Further, models with fine meshes and many mesh points can take a long time to visualise. This is particularly true of three-dimensional models, where it is often necessary to identify which mesh points lie on the outer surface of the model before visualisation can proceed. In order to produce usable images as quickly as possible, some packages use routines to produce a close approximation to the mesh that can be visualised more quickly than the full mesh.

The algorithms used for this procedure require a tolerance in order to decide whether or not a mesh point can be ignored. If the distortion to the geometry visualisation caused by omitting the mesh point from the visualisation is less than the tolerance, the mesh point is ignored. The user can change this tolerance. It is likely that the default tolerance will be suitable for the majority of applications, but models in which the details of the geometry are crucial may require the tolerance to be altered so that only a very low level of mesh distortion is permitted.

6.4 Uncertainties and continuous modelling

The importance of visualisation of uncertainties to metrology has been discussed in Chapter 5. The contents of that Chapter apply as much to continuous models as to any other metrology visualisations. The visualisation of uncertainties in continuous modelling is crucial if users are to have confidence in the model results. It is important for users to understand the uncertainties associated with the results prior to making decisions based on visualisations, particularly since visualisations can encourage users to forget that the quantities they are viewing have uncertainties associated with them. Use of visualisations of the model results and the associated uncertainties will improve the user's understanding of the overall behaviour of the model results. However, since continuous model results are frequently two- or three-dimensional, the visualisation of their uncertainties can be problematic.

When visualising uncertainties, as explained in Chapter 5, it is necessary for the user to consider the two questions:

- Which output quantities are of interest?
- How are the uncertainties to be represented?

The answers to these two questions are often inter-dependent. For example, representing the uncertainty as a distribution function may limit the range of output quantities that can be considered. Similarly, if the full results of a three-dimensional model are required, it may only be possible to visualise expectations and the associated standard uncertainties, since typically the results of such models are given at several thousand spatial points and so visualisation of correlations or distribution functions is not feasible.

6.4.1 Choice of output quantity

The most obvious choice of output quantity is the complete set of model results. However, the majority of continuous models produce results at several thousand points, and so the retention and manipulation of distribution function information relating to all these points can often be prohibitive. Additionally, some models generate multiple results at each point, not all of which are necessarily of interest.

In many cases, the user is only interested in a subset of the results. For example, measurements are often only taken at a small number of points, and the user may be particularly interested in the uncertainties associated with the results at those points. Similarly, the user may be most interested in the results on the boundaries of the model. If a smaller number of points is of interest, manipulating and visualising the distribution functions becomes less problematic. Reduction of the results of interest to a one- or two-dimensional set can increase the range of visualisation methods that can be used, and often makes the visualisation simpler and easier to interpret.

The user may be interested in some key feature of the results such as the location or value of the maximum and minimum of the results. There may be a single quantity that can be derived from the results that is of particular interest, such as an elastic modulus or a thermal conductivity. In such cases, the visualisation can proceed in the same way as the

visualisation of the results of a discrete model, since the origin of the quantities is unimportant (Chapter 5).

6.4.2 Choice of representation of the uncertainty

As mentioned in Chapter 5, the probability distribution of the value of a quantity can be represented in several ways. A plot of the distribution function or probability density function will give the most information, but often it is more convenient to summarise the quantity as an expectation, standard deviation or standard uncertainty, and a coverage interval. However, if the quantity does not follow one of the standard distributions, a significant amount of information can be lost if a distribution function is summarised in this way.

The representation of uncertainties associated with two- and three-dimensional results can be problematic. If the distribution functions associated with all the results are displayed simultaneously, the visualisation can be cluttered and confused. If the results are displayed as expectations and standard uncertainties, the visualisation will be clear but detailed information will be lost.

Many of the most common discretisation methods used for continuous models produce large linear systems of equations that link the values of the variables at the mesh points. If the model is time-dependent, the values of the variables at each new time step are usually written as a linear combination of the values at the previous time step. This strong inter-dependency between variable values means that the model results at different mesh points are typically mutually dependent. This correlation between results further complicates the visualisation of the uncertainties.

6.4.3 Software for visualising uncertainties associated with continuous models

There are fewer packages suitable for visualisation of uncertainties associated with continuous models than there are for visualisation of continuous model results. Many continuous modelling software packages consist of a pre-processor for model development, a program for solution of models, and a post-processor for output and visualisation of results. In general, the post-processing parts of these packages will only display the results calculated by the solution program: they are not usually sufficiently flexible to display quantities, such as uncertainties, that are calculated by other means.

If the results of a three-dimensional model have been chosen to be the quantities of interest, the choice of software packages may be quite restricted, particularly if the model geometry is complicated. Visualisation of such results requires software to be able to read in and interpret nodal positions and element connectivities, and to be able to display the uncertainties on a visualisation of the problem geometry. Some professional visualisation packages, such as IRIS Explorer (<http://www.nag.co.uk/Welcome.IEC.html>) and Amira (<http://www.amiravis.com/>) are able to read and interpret geometric output from common computer aided design (CAD), computational fluid dynamics (CFD) and finite element (FE) packages such as IDEAS, Phoenix, and Hypermesh. Since these packages also accept text file input, it is possible to read a complicated geometry from the CAD or

FE package output files, read the associated uncertainties from a text file, and visualise the uncertainties that way.

For simpler geometries, the range of packages able to visualise results may be larger as it will be more straightforward to define the problem geometry within the visualisation package itself.

6.5 Recommendations

- Check the interpolation, extrapolation and smoothing components of the visualisation software:
 - If the derivation of the discretised problem involved approximation functions such as polynomials, visualise the results using the same type of functions.
 - Visualise derivative quantities, such as stresses and strains, using the derivatives of the approximation functions.
 - If the derivation of the discretised version of the problem did not involve approximation functions, assume that the results are constant within an appropriate visualisation element.
 - Wherever possible, turn smoothing effects off to check that they are not disguising unexpected discontinuities in the results.
 - Compare visualisations using smoothed and unsmoothed results to increase confidence in the model results. Visualisations of a well-validated model with continuous results with and without smoothing should look broadly the same.
- Check the visualisation parameter settings:
 - Read the package documentation to find out which parameters affecting the display have default values.
 - Consider which values of these parameters are most suitable for the results being visualised.
 - If it is not clear which value is most suitable, adjust the value and rerun the visualisation to obtain an idea of the effect of the parameter.
 - Consider limiting the visualisation to a small part of the model if displaying the full model would produce a misleading picture of the results.
- Choose a sensible scale for displaying results.
- Check tolerances used in algorithms for displaying results.
- Visualisation of the uncertainties associated with the results of a continuous model can increase the user's understanding of the results and can improve confidence in any decisions that are made based on visualisations.

Chapter 7

Examples

This chapter includes examples concerned with:

1. Visualising models of experimental data (section 7.1)
2. Visualising uncertainties associated with the calibration of microwave power meters (section 7.2)
3. Visualising measurements of a sphere (section 7.3)
4. Visualising the results of continuous modelling (section 7.4)
5. Visualising ordinal and categorical data (section 7.5)
6. Validating a visualisation (section 7.6).

7.1 Visualising models of experimental data

7.1.1 Introduction

Figure 7.1 shows a data set obtained from an experiment to characterise the motion of a piezoelectric transducer using a laser-interferometer measurement system [13]. The data set comprises $m = 100$ points (x_i, y_i) , $i = 1, \dots, m$, and it is required to approximate the data by a continuous curve. In the absence of a model for the data based on physical considerations, and because the underlying curve possesses a single turning point and exhibits no rapid change in curvature, it is appropriate to model the data by a polynomial. Accordingly, the class of polynomials is considered as a suitable family of models for the data. The family has the property that each member ‘contains’ all previous members, i.e., polynomials of order n (degree $n - 1$) ‘contain’ polynomials of all lower orders.

Polynomials of orders $n = 1, 2, 3, \dots$ (degrees $0, 1, 2, \dots$) are fitted to the data by the method of least-squares. It is assumed that the measurements x_i of the independent variable are accurate compared to those y_i of the dependent variable. Furthermore, all

data points are accorded the same weight in the least-squares fitting, because there is no reason to expect that some measurements y_i are any more accurate than others. For reasons of numerical stability, the polynomial $p_n(x)$ of order n (degree $n - 1$) is represented in terms of Chebyshev polynomials in a variable X normalised to the interval $[-1, +1]$ [15].

The objective is to use visualisation to help make an informed choice of the order of the polynomial fit to be used to approximate the data. This example shows the value in using simple plots to support the understanding of experimental data as well as the understanding of data modelling.

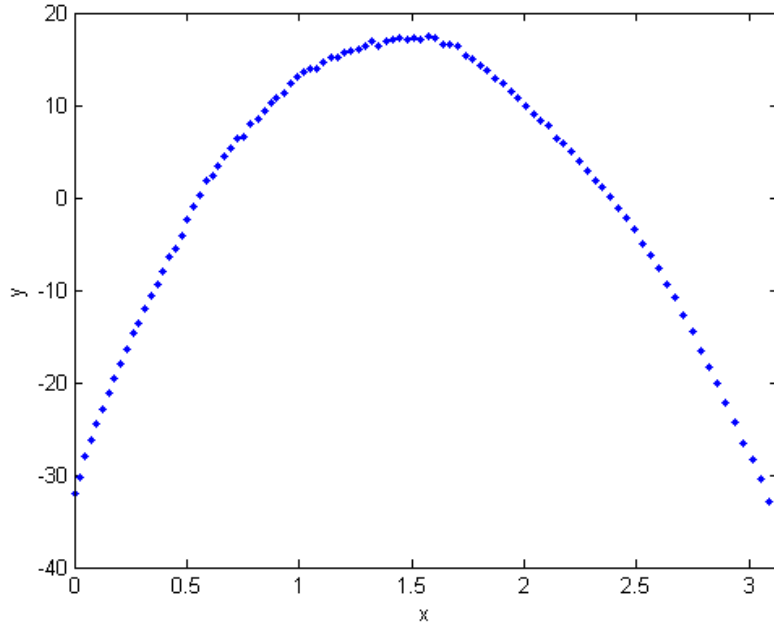


Figure 7.1: Experimental data obtained from an experiment to characterise the motion of a piezoelectric transducer.

7.1.2 Visualisations

For each polynomial fit $p_n(x)$, the root mean square residual deviation,

$$s = \sqrt{\frac{\sum_{i=1}^m (y_i - p_n(x_i))^2}{m - n}},$$

provides a measure of the departure of the fit from the data. The measure $(m - n)s^2$, regarded as a function of polynomial order n , will be non-increasing, because $p_n(x)$ is the polynomial determined to minimise $(m - n)s^2$ and polynomials form a family of models (as described above). It is expected that the measure s will generally decrease, but will stabilise to an essentially constant value. The value provides a measure of the accuracy of the measurements. In Figure 7.2 values of the root mean square residual deviation s are

plotted against polynomial order. There is an appreciable drop in the value of s between orders two and three. This outcome is consistent with the fact that visually the shape of the underlying curve is approximately quadratic. It is difficult to use this graph, however, to see how s is changing for orders greater than four.

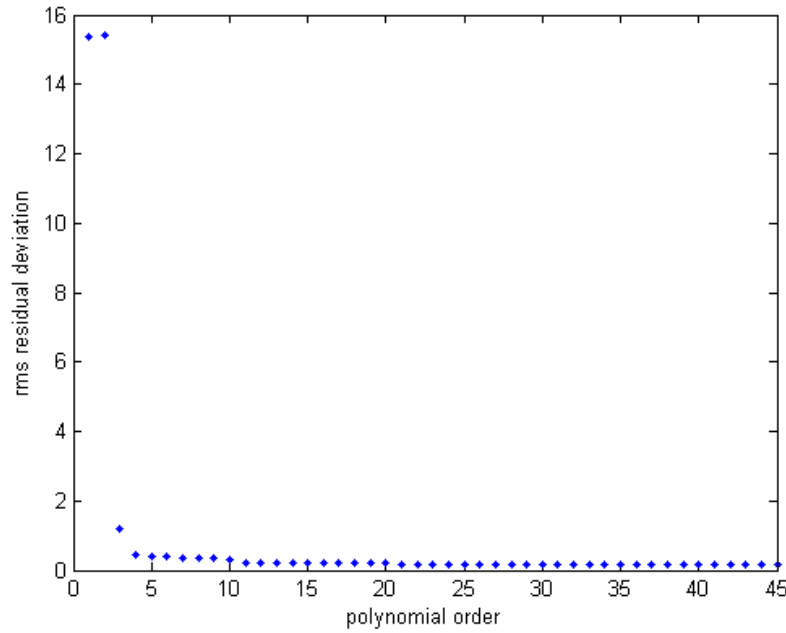


Figure 7.2: Root mean square (rms) residual deviation plotted against the order of polynomial fit to the data shown in Figure 7.1.

Figure 7.3 displays the same information as Figure 7.2, but with the logarithm (to base ten) of the root mean square residual deviation plotted against polynomial order. The transformation of the data helps emphasise the differences between the values of s when those values are small. The graph indicates that the value of s stabilises when the polynomial order is (about) 11, and again when the polynomial order is (about) 23. An examination of the numerical values of s establishes for that for the range of polynomial orders shown, s takes its minimum value when the polynomial order is 28.

Another measure of the departure of a fit from the data is provided by the maximum absolute residual deviation,

$$r = \max_{i=1,\dots,m} |y_i - p_n(x_i)|.$$

Figure 7.4 displays the logarithm of the maximum absolute residual deviation against polynomial order. There is an appreciable drop in the value of r between orders two and three (as in Figure 7.3) and, thereafter, the value of r continues to generally decrease.

To explore the ‘trade-off’ between root mean square residual deviation and maximum absolute residual deviation, we can plot as in Figure 7.5 values of s and against the corresponding values of r . The graph confirms that the polynomial fit of order 28 gives the smallest value of s , and shows that the fit of order 23 is almost as good in terms of root

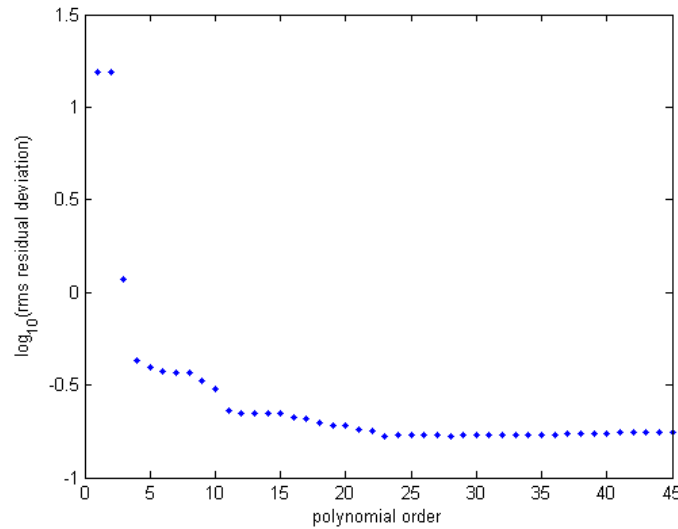


Figure 7.3: As Figure 7.2, but plotting the logarithm (to base ten) of the residual mean square residual deviation.

mean square residual deviation but is poorer in terms of maximum absolute residual deviation. Two graphs are presented: one to give a general impression of the data and set the interesting data in context, and the other to focus on the interesting data (that for polynomials of order 20 and above).

It is also useful to examine, for each polynomial order, the residual deviations,

$$e_i = y_i - p_n(x_i), \quad i = 1, \dots, m,$$

themselves. Figure 7.6 shows the residual deviations, mapped to a colour scale, plotted against the data x -values (on the horizontal axis) and the polynomial order (on the vertical axis). In this figure a colour scale is chosen to give a reasonable balance in the impression of large positive and large negative residual deviations and to avoid relying on the distinction between red and green (which can be difficult for those with colour blindness). The figure shows that for low polynomial orders (less than ten) the residual deviations of large absolute value tend to be grouped together. As the polynomial order is increased, the grouping becomes much less pronounced, and the range of values of the residual deviations also decreases. Figure 7.7 shows the same information as in Figure 7.6, but using the default colour scale provided by MATLAB. It is perhaps more difficult to understand the figure and to draw conclusions from it.

To compare two polynomial fits it can be useful to examine directly the residual deviations for the two fits, as in Figures 7.8 and 7.9. Figure 7.8 shows the residual deviations associated with polynomial fits of orders 10 and 11, which correspond to where the root mean square residual deviation first appears to stabilise (Figure 7.3). The residual deviations associated with the two fits are plotted with the same vertical scales, to aid comparison, and the experimental data is also shown, in order to put the magnitude of the deviations in context. The deviations associated with the fit of order 11 are noticeably more ‘random’ than those associated with the fit of lower order: this is particularly so for

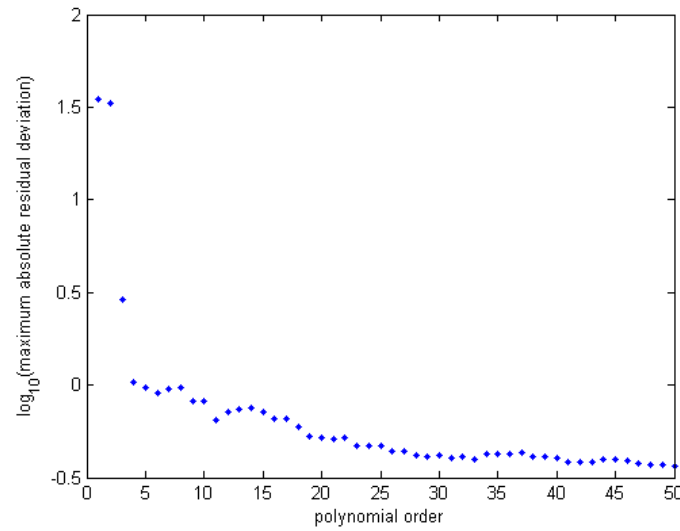
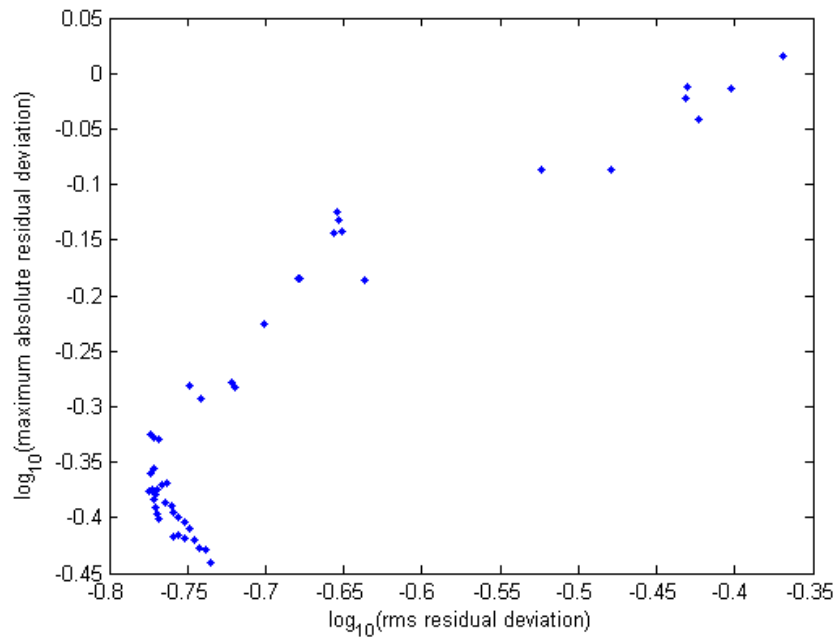


Figure 7.4: The maximum absolute residual deviation plotted against the order of polynomial fit.

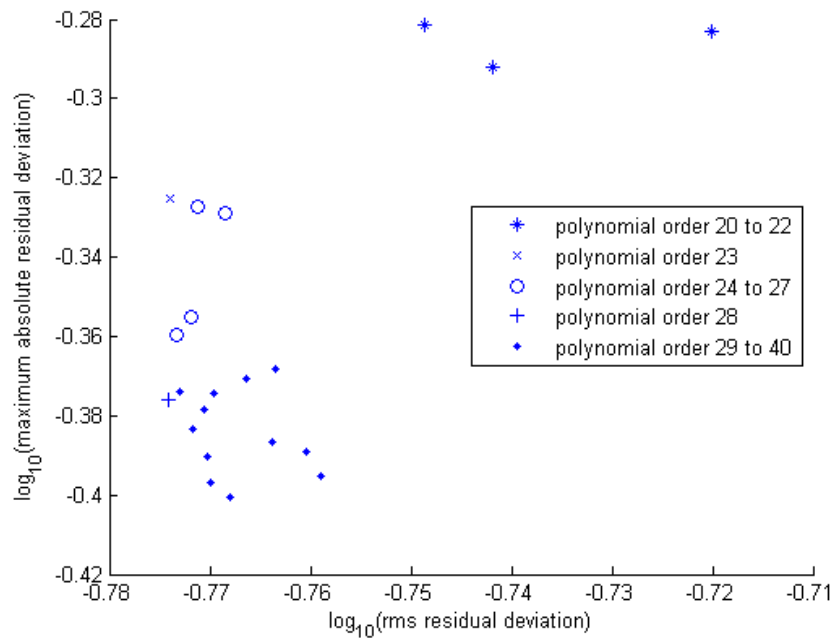
$x > 2$. Figure 7.9 shows, similarly, but with a smaller vertical scale, the residual deviations associated with polynomial fits of orders 23 and 28. The former marks the start of the second range of polynomial orders for which the root mean square residual deviation has stabilised. The latter corresponds to the smallest root mean square residual deviation within that range. There appears to be much less change in the ‘nature’ of the residual deviations between orders 23 and 28 (compared to the change between orders 10 and 11). It would appear, on this basis, that there is little benefit in increasing the order of the polynomial fit beyond 23.

The residual deviations provide information about a polynomial fit at the points at which experimental data are available. It is also important to examine the behaviour of the fit between the data points as a continuous function of the independent variable x . A plot of the data together with the polynomial fit can be difficult to interpret because the differences between the data and fit can be small compared to the range of data values. One approach is to plot the data with the polynomial fit after subtracting a low order polynomial from the data and the fit.

The approach is illustrated by the graphs shown in Figure 7.10. Graph (a) in the figure shows the experimental data with the polynomial fit of order four. Graph (b) shows the differences between the data and the polynomial fit of order four (shown as data), and the difference between the polynomial fits of orders 4 and 11 (shown as a curve). Graphs (c) and (d) present similar information, but for polynomial orders of 23 and 28. The graphs show that the polynomial fits of orders 23 and 28 exhibit some spurious behaviour (particularly near to the start and finish of the data) that may make the fits less suitable for interpolating between the data points in these regions. Conversely, the polynomial fit of order 11 is ‘smoother’, but does not describe all the features represented by the data.



(a)



(b)

Figure 7.5: The maximum absolute residual deviation plotted against the root mean square residual deviation. The lower graph shows the interesting detail associated with the upper graph.

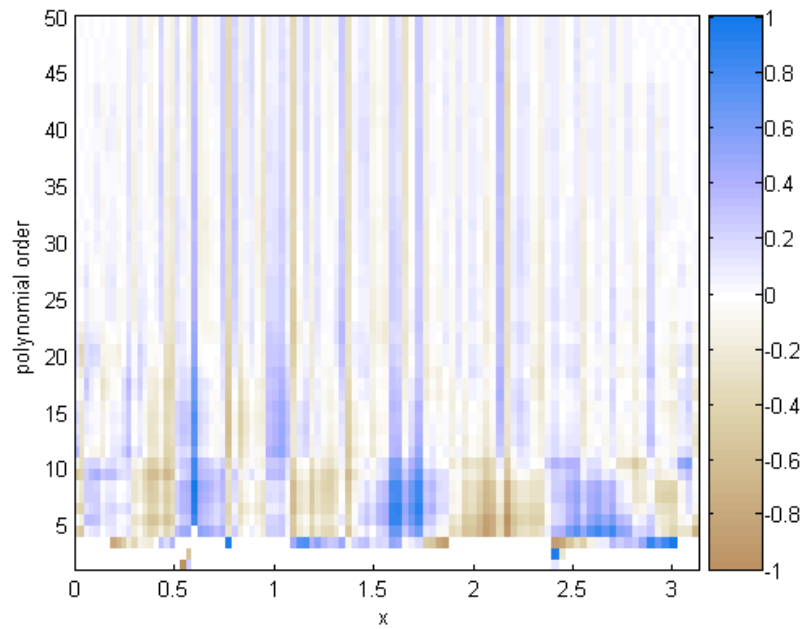


Figure 7.6: Residual deviations plotted against data x -value (horizontal axis) and order of polynomial fit (vertical axis).

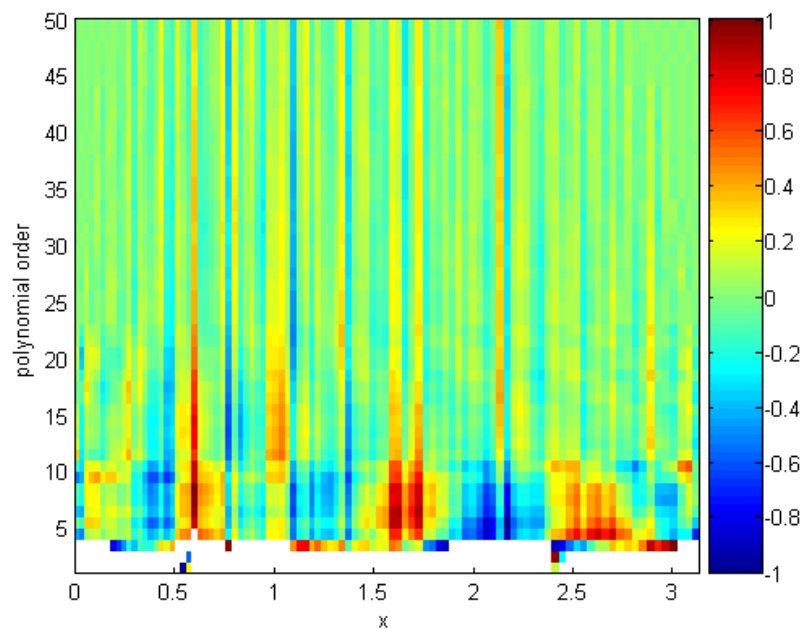


Figure 7.7: As Figure 7.6, but using the default colour scale.

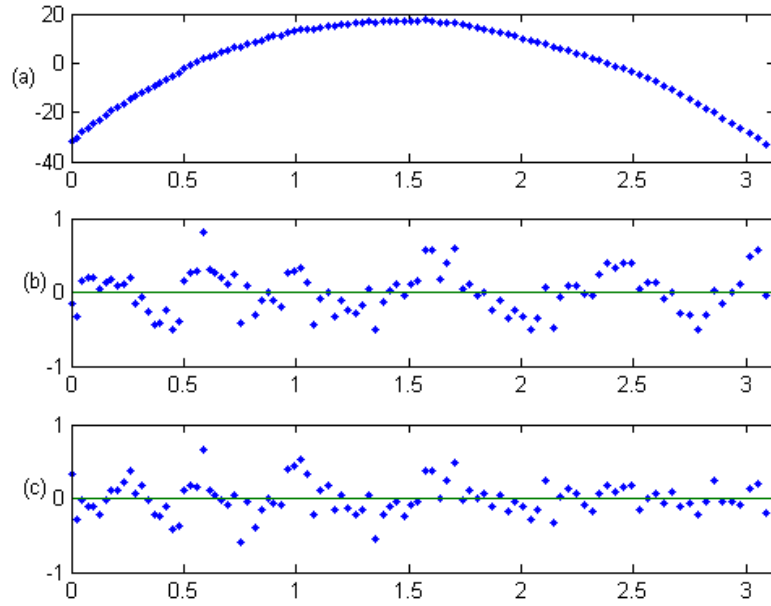


Figure 7.8: (a) Experimental data. (b) Residual deviations associated with the polynomial fit of order 10. (c) As (b), but for the polynomial fit of order 11.

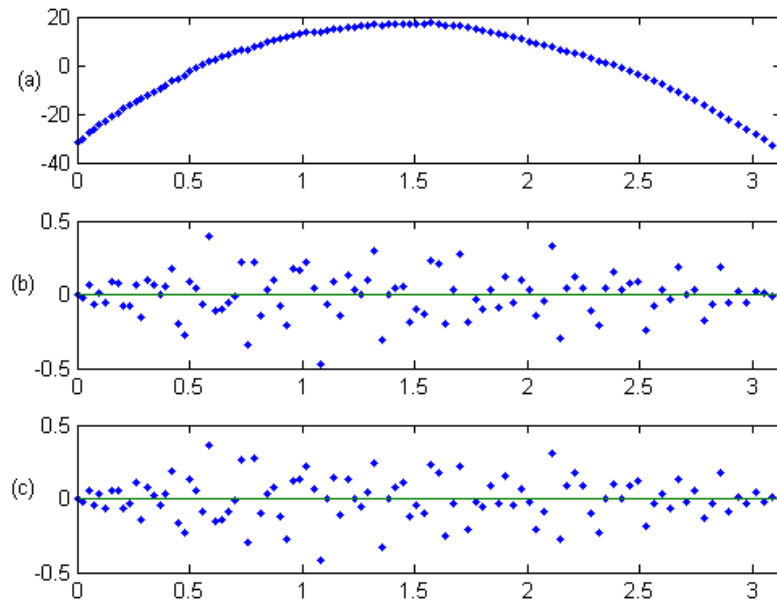


Figure 7.9: (a) Experimental data. (b) Residual deviations associated with the polynomial fit of order 23. (c) As (b), but for the polynomial fit of order 28.

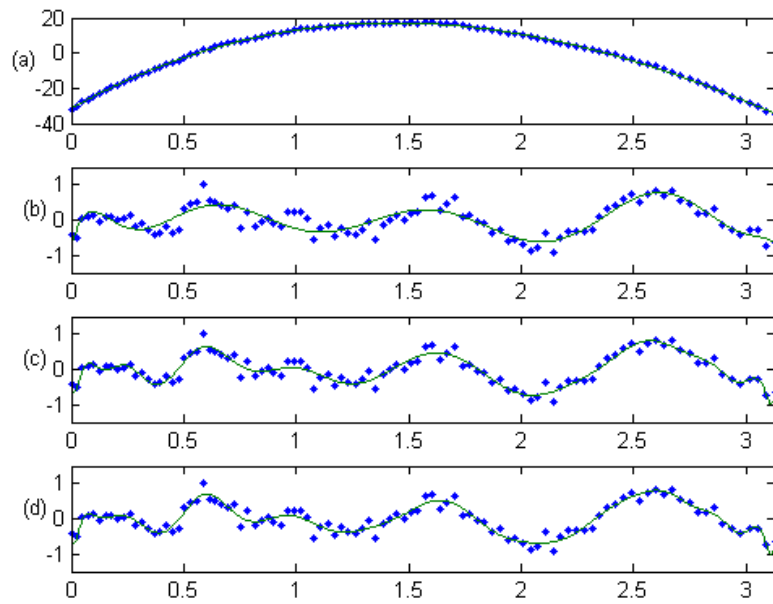


Figure 7.10: (a) Experimental data plotted with the polynomial fit of order 4, and residual deviations associated with the polynomial fit of order 4 plotted with the difference between the polynomial fits of (b) orders 4 and 11, (c) orders 4 and 23, and (d) orders 4 and 28.

7.1.3 Summary

All the visualisations considered in this example (perhaps with the exceptions of Figures 7.6 and 7.7) are straightforward to produce. The example has shown how simple plots can be used to support the understanding and modelling of experimental data.

7.2 Visualising uncertainties associated with the calibration of microwave power meters

7.2.1 Introduction

During the calibration of a microwave power meter, the power meter itself and a standard power meter are connected in turn to a stable signal generator. The power absorbed by each power meter will in general be different because their input voltage reflection coefficients are not identical. The ratio M of the power P_M absorbed by the power meter being calibrated and that, P_S , absorbed by the standard power meter is [26, 27]

$$M = \frac{P_M}{P_S} = \frac{1 - |\Gamma_M|^2}{1 - |\Gamma_S|^2} \times \frac{|1 - \Gamma_S \Gamma_G|^2}{|1 - \Gamma_M \Gamma_G|^2}, \quad (7.1)$$

where Γ_G is the voltage reflection coefficient of the signal generator, Γ_M that of the power meter being calibrated and Γ_S that of the standard power meter. The voltage reflection coefficients are complex-valued as well as, in general, being functions of frequency f . This power ratio is an instance of ‘comparison loss’ [3, 19].

Consider the simple case where both the standard and the signal generator are reflectionless, i.e., $\Gamma_S = \Gamma_G = 0$, and where measurements are made of the real and imaginary parts, viz., X and Y , of $\Gamma_M = X + jY$, where $j = \sqrt{-1}$. Since $|\Gamma_M|^2 = X^2 + Y^2$, the comparison loss in formula (7.1) becomes

$$M = 1 - |\Gamma_M|^2 = 1 - X^2 - Y^2. \quad (7.2)$$

In order to determine whether a power meter has met a specification value for its comparison loss, it is necessary to obtain a measurement of Γ_M and the associated uncertainty. The measurement is expressed in terms of estimates x and y of the real and imaginary parts X and Y of Γ_M , and the uncertainty in terms of standard uncertainties $u(x)$ and $u(y)$ and correlation coefficient $\rho(x, y)$ associated with the estimates. The operator testing the meter (or the designer of the equipment used to perform the test) is interested in knowing, for a given measurement of Γ_M , over what range of values the uncertainty components $u(x)$, $u(y)$ and $\rho(x, y)$ can vary for the measured comparison loss to be within a given specification limit. Therefore, for a given measurement of Γ_M , the following steps are undertaken:

1. Determine an estimate m of the value of the comparison loss M (from the measurement of Γ_M).
2. Check whether m lies ‘within’ (i.e., is greater than) the comparison loss specification limit M_{spec} , i.e., whether $m > M_{\text{spec}}$.

3. If the answer is affirmative, determine the amount by which m exceeds M_{spec} to establish an upper limit for the *expanded* uncertainty $U(m)$ that can be tolerated in the assessment of m against the specification, i.e.,

$$U(m) \leq m - M_{\text{spec}}. \quad (7.3)$$

This situation occurs, for example, when a designer of a product's test bench needs to establish the level of uncertainty that can be tolerated by a given test in order that a certain type of product can be tested in an efficient manner. This decision can affect both the choice of test type and test equipment used to ensure adequate quality control (i.e., including both economic and scientific factors in the decision concerning the fitness for purpose of the test). The symbol U represents an expanded uncertainty corresponding to a suitable coverage probability p . For most cases of compliance assessment against a specification limit (in this area of testing), p is often taken as 0.95. (This form of compliance assessment is often called 'inset limits'.)

4. If the answer is negative, regard the meter as unsatisfactory, since its performance is inconsistent with the specification limit.

The aim of the visualisation is to assist the designer to identify those combinations of the components of the uncertainty associated with a measurement of Γ_M that can be accommodated in keeping with inequality (7.3). In this example, the 'independent variables' for a visualisation are $u(x)$, $u(y)$ and $\rho(x, y)$, and the 'dependent variable' $u(m)$. Visualisations are provided of *isosurfaces* composed of those points $(u(x), u(y), \rho(x, y))^T$ that correspond to constant values of $u(m)$.

A treatment of the evaluation of uncertainty for the model (7.2) of comparison loss by applying the law of propagation of uncertainty and the propagation of distributions implemented using a Monte Carlo method is available [9], together with information about the calculation and display of the above isosurfaces.

7.2.2 Visualisations

A web-based interactive visualisation of the isosurfaces obtained using the law of propagation of uncertainty (LPU) and a Monte Carlo simulation (MCS) is available in the electronic form of this document. Henceforth, the isosurface obtained using LPU will be referred to as the 'LPU surface' and that obtained using MCS as the 'MCS surface'. Information concerning its use is given below.

Figure 7.11 illustrates an example of an isosurface generated by this interactive visualisation for a typical range of variation in $u(x)$ and $u(y)$. Since establishing (even a coarse approximation to) the correlation coefficient $\rho(x, y)$ associated with x and y is (in practice) very difficult, the vertical axis in the visualisation is used to represent the full range $[-1, 1]$ for this quantity.¹

¹The label ' $r(x, y)$ ' is used on the vertical axis (in place of ' $\rho(x, y)$ ') in the visualisation tool.

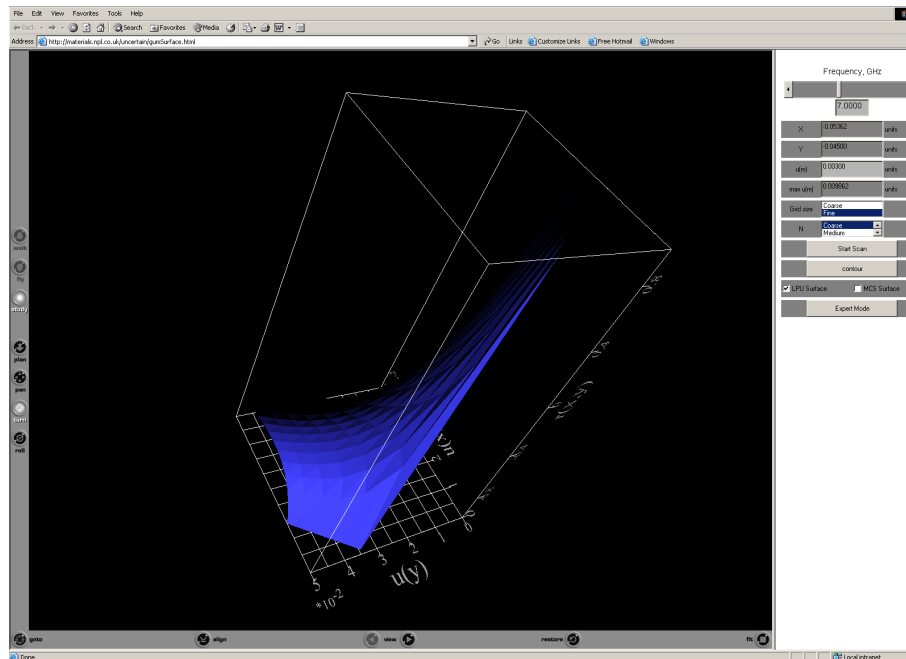


Figure 7.11: Illustrative surface and controls for a surface of constant standard uncertainty $u(m)$ associated with comparison loss.

Operating the visualisation

There are two sets of controls available for operating the visualisation. The first set, displayed on the left-hand side of the screen (Figure 7.11), comes from the VRML browser and enables the user to move the scene on the computer screen. These controls are labelled *Plan*, *Pan*, *Turn* and *Roll*, in accordance with the type of movement required by the user. Each control is activated by selecting (i.e., ‘clicking’) the appropriate on-screen icon. The scene is then moved, using the selected control, by the user clicking and dragging the scene across the active area of the visualisation.

The second set of controls, displayed on the right hand side of the screen (see Figure 7.12 for a magnified image), comes from the application and enables the user to alter the parameters used to generate the visualisation. These controls are labelled *Frequency*, *X*, *Y*, *u(m)*, *max u(m)*, *Grid size* and *N*. In addition, there are ‘buttons’ labelled *Start Scan*, *contour*, *LPU surface*, *MCS surface* and *Expert mode*. Each of these controls and buttons is explained below.

Controls

Frequency. A slider allowing frequencies to be accessed at integer GHz values ranging from 0 to 18. Alternatively, any frequency value can be entered manually into the box immediately below the slider.

X and Y. Values determined automatically, based on the selected frequency and a

Frequency, GHz

7.0000

X	-0.05362	units
Y	-0.04500	units
u(m)	0.00300	units
max u(m)	0.009862	units
Grid size	Coarse Fine	
N	Coarse Medium	
Start Scan		
contour		
☒ LPU Surface ☐ MCS Surface		
Expert Mode		

Figure 7.12: Magnification of the controls in Figure 7.11.

‘spiral’ model for the estimates x and y [9]. They are used to simulate the behaviour as a function of frequency of a typical power meter. The use of simulated values for X and Y is adequate for the purposes of this ‘demonstrator model’ of the visualisation software—a ‘production version’ would be modified to import values of X and Y direct from measurements made using laboratory instruments such as Vector Network Analyzers.

$u(m)$. Required value of $u(m)$. The isosurface will only appear on the computer screen if $u(m) < \max u(m)$, otherwise a warning message is displayed.

$\max u(m)$. Maximum calculated value for $u(m)$ occurring within the scale ranges set for the visualisation, i.e., the displayed ranges for $u(x)$ and $u(y)$. It is determined automatically.

Grid size. *Coarse* or *Medium*—*Coarse* being the default setting. It determines the number of points on a grid within the visualisation ‘space’ where uncertainty evaluations are performed. *Medium* will generally produce smoother surfaces, but will increase the time taken to compute the surfaces.

N . *Coarse*, *Medium* or *Fine*, corresponding, respectively, to $N = 1\,000$, $10\,000$ and $100\,000$. It determines the number N of Monte Carlo trials used to perform the Monte Carlo calculation. *Fine* should be used to obtain the most accurate results, although the time taken to compute the MCS surface will be greater. *Coarse* or *Medium* settings are recommended when exploring general features in the computed surfaces.

Switches

Start Scan. Allows the frequency to be ‘swept’ automatically, in real time, in integer GHz values, from 1 to 18. Only the ‘LPU surface’ is displayed, since the computation of the MCS surface can take considerable time. The scan can be stopped at any time by pressing the *Cancel* button, which replaces the *Start Scan* button during the scanning process.

Contour. Presents a series of ‘slices’ through the isosurfaces at values of correlation coefficient, $\rho(x, y)$, corresponding to -1.0 , -0.5 , 0.0 , $+0.5$ and $+1.0$. Each slice shows a two-dimensional grid, with axes $u(x)$ and $u(y)$, with contour lines for different values of $u(m)$. The numbers on these contour lines correspond to the scale of values for $u(m)$, from zero to $\max u(m)$, given at the bottom of the screen. An additional contour line, shown in green, is the contour line corresponding to the value for $u(m)$ selected by the user and displayed as the surface in the main visualisation window. Slices through both LPU and MCS surfaces can be displayed using the *Contour* option. However, to display slices through the MCS surface, the MCS surface button must first be pressed (below).

LPU surface. A switch allowing the LPU surface to be turned on or off, i.e., displayed or not displayed in the visualisation.

MCS surface. A switch allowing the MCS surface to be turned on or off. Note that the MCS surface must be re-calculated and re-displayed each time there is a change in the values in the parameter list, e.g., *frequency*, *Grid size* or N , but not $u(m)$. To

indicate this requirement, a dialogue box, containing the message ‘Update MCS surface—you need to recalculate the MCS data’, appears automatically whenever there has been a change in the parameter list settings. A re-calculation is initiated by clicking the ‘continue’ button in this dialogue box. This aspect is most important, since any MCS surface that is already shown, but has not been re-calculated and re-displayed, is likely to be based on previous parameter value settings and will therefore, in general, not be the correct surface for the modified display.

Expert mode. Provides access to a larger set of controls and buttons intended for diagnostic and training purposes.

Example

Suppose the specification limit for comparison loss, at 12.247 GHz, is given as 97.5 %, i.e., $M_{\text{spec}} = 0.975$. This is the smallest value of comparison loss for the meter to be within its specification.² During the test measurements on the meter, the real and imaginary parts of Γ_M were found to be -0.052 and 0.111 , respectively, and hence $|\Gamma_M|^2 = 0.015$. The comparison loss is therefore measured as

$$m = 1 - |\Gamma_M|^2 = 0.985.$$

The expanded uncertainty $U(m)$ associated with the comparison loss measurement must conform to

$$U(m) \leq m - M_{\text{spec}} = 0.985 - 0.975 = 0.010$$

to ensure that this meter falls within specification. The visualisation (Figure 7.13) is used to assist in identifying those combinations of $u(x)$, $u(y)$ and $\rho(x, y)$ that satisfy the condition $u(m) = 0.005$. Assuming a Gaussian distribution is assigned to the value of M , and p , the coverage probability, is 0.95, it follows that $U(m) = k_{0.95}u(m)$, where $k_{0.95} \approx 2.0$ [5]. For the purposes of this ‘demonstrator model’ of the visualisation software, the standard uncertainty $u(m)$, rather than the expanded uncertainty $U(m)$, is visualised. A production version of the software would provide a visualisation of $U(m)$, which is the more useful quantity for verifying the performance of a power meter.

These combinations can be seen by inspecting the contour ‘slices’ through the visualisation surface, as can general trends indicated by the surface as a whole. It is important to examine the entire surface for the presence of any unusual behaviour that may be present at values of $\rho(x, y)$ between the contour ‘slices’. Such behaviour includes cusps, singularities, inflections and any other rapid variations in the shape of the surface as a function of $\rho(x, y)$. If the choices are limited to instances where $u(x) = u(y)$, it is appropriate to see where the straight line $u(x) = u(y)$ intersects the contour $u(m) = 0.005$.³

Now, if, for example, $u(x) = u(y) = 0.02$ can easily be achieved, but no indication of the value of $\rho(x, y)$ is available, the following situations can be seen from the contour slices:

²During this example, for simplicity, only the ‘LPU surface’ is considered.

³This is a realistic limitation in the uncertainty structure. It is used currently in PIMMS [25], NPL’s Primary Impedance Microwave Measurement System.

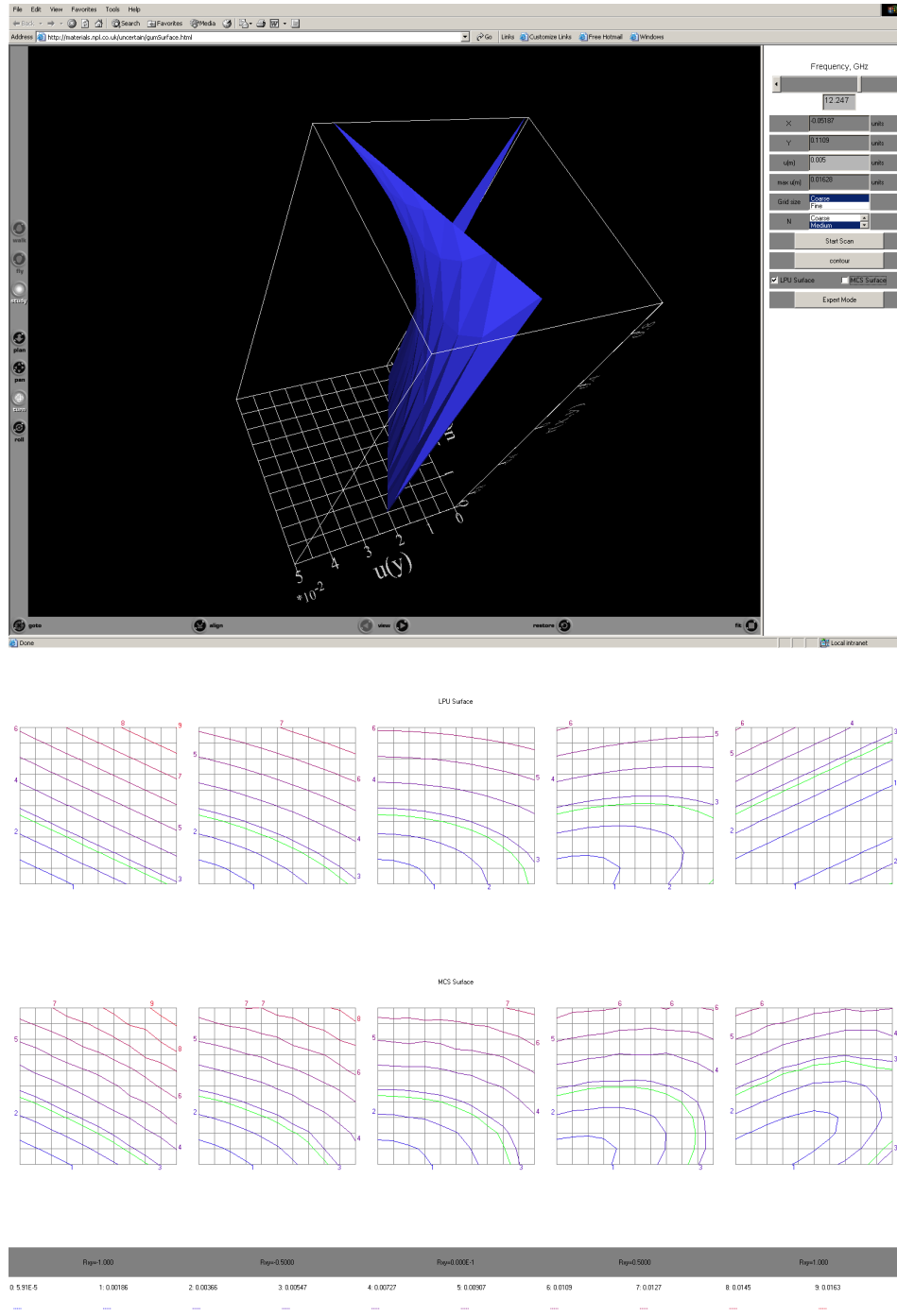


Figure 7.13: The LPU surface corresponding to a frequency of 12.247 GHz and $u(m) = 0.005$, and (bottom) the contours defining the LPU surface and the corresponding MCS surface. The MCS surface is obtained using the ‘medium’ setting of $N = 10\,000$ trials.

$\rho(x, y)$	$u(m)$	The product
-1	> 0.005	Probably rejected as non-compliant
-0.5	> 0.005	Probably rejected as non-compliant
0	≈ 0.005	Compliance cannot be clearly established
0.5	< 0.005	Can be accepted as compliant
1	< 0.005	Can be accepted as compliant

In order to demonstrate satisfactory compliance with the specification, regardless of the value of $\rho(x, y)$, the uncertainty associated with Γ_M must be reduced so that $u(x) = u(y) < 0.015$. Such a reduction could be achieved, for example, by using better test equipment to perform the measurements, or by carrying out investigative work to reduce key components in the equipment's uncertainty budget.

7.2.3 Summary

This example has shown how visualisation can be used to inform an operator about the conditions of measurement necessary to ensure the compliance of a product. It has illustrated how different visualisation tools, including plots of isosurfaces and contours, as well as interactive visualisation, can be used to explore a 'mathematical' relationship between more than three variables (here, $f, u(x), u(y), \rho(x, y)$ and $u(m)$) and, in particular, how one of those variables (here, $u(m)$) depends on the others.

7.3 Visualising measurements of a sphere

7.3.1 Introduction

The data set for this visualisation arises from optical measurements of a precisely manufactured glass sphere with a nominal diameter of 12.5 mm. The optical system measures the distance between the glass surface and a reference spherical wavefront. The reference wave front subtends an angle of 84 degrees at the centre of the sphere, and thirty-six sets of measurements are taken with the sphere rotated in ten degree steps around a common axis through the ‘poles’ of the sphere. The data at each position is provided as a measurement (of distance) between the sphere and the reference wavefront together with spherical polar co-ordinates giving the position of the measurement on the reference wavefront.

Examination of the spherical polar co-ordinates for the measurements suggests that the data is sampled on a rectangular grid in a plane with the projection between the reference wavefront and the plane as shown in Figure 7.14. The objective is to use visualisations of the data to help verify, or otherwise, the assumption made about the data sampling.

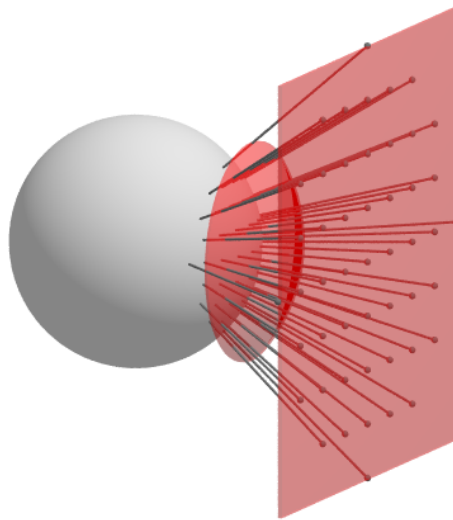


Figure 7.14: Measurements of the sphere are assumed to taken on a rectangular grid projected from a plane to the reference wavefront.

7.3.2 Visualisations

In order to combine the thirty-six sets of measurements each one is projected onto the surface of a cylinder aligned with the axis of rotation of the sphere. The cylinder provides a common co-ordinate system for viewing the data. Three of the projected images are shown in Figure 7.15. At first sight the projected images seem consistent because there are features within the data which appear to remain stationary as the area of measurement is

moved. To check this more carefully, the sequence of images can be viewed as an animation. The human visual system is extremely sensitive to movement, and a standard method for comparing two very similar images is to rapidly switch between them. An animation shows that the features in the data do not stay completely stationary. The effect is too small to be seen in Figure 7.15, but is significant for the application.

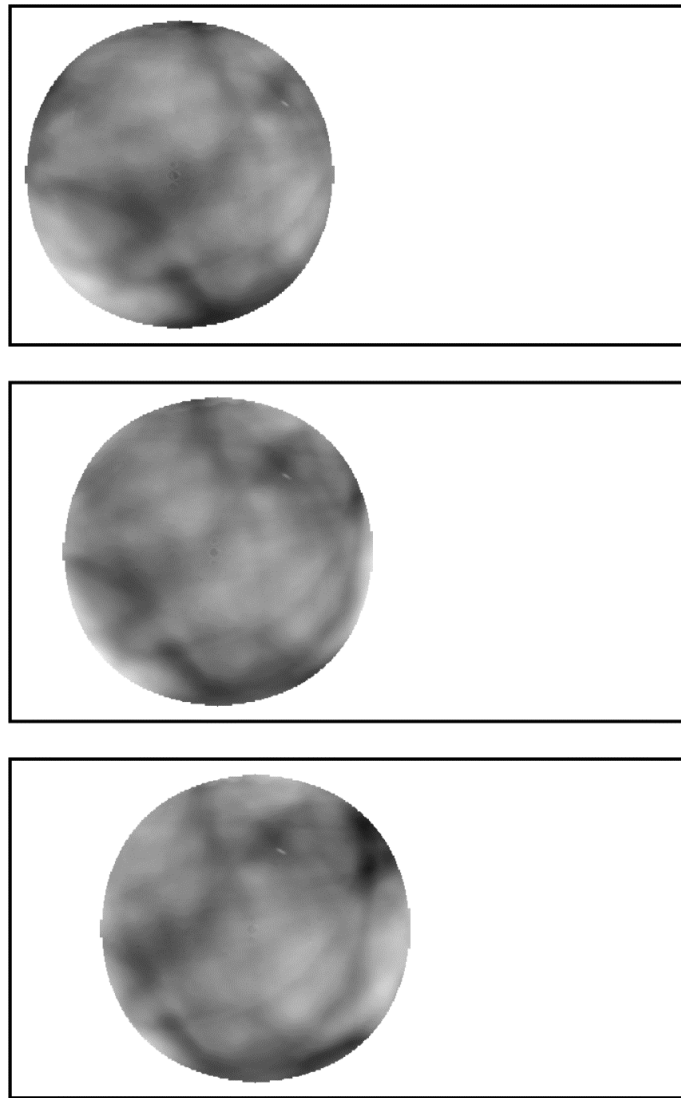


Figure 7.15: Measurements of the sphere at ten degree increments in rotation viewed in a common co-ordinate system.

A radial distortion applied to the data is sufficient to bring the features within the individual images into co-incidence. The distortion would be consistent with the optical system used to generate the wave front not having a flat principle plane. Without the use of visualisation, the oversight in the communication of information from the optical designer to those analysing the experimental data would not have been identified.

For data corrected for radial distortion, it is possible to combine the individual images of the surface into a single image, as in Figure 7.16. The data can also be visualised as a three-dimensional projection as shown in Figure 7.17.

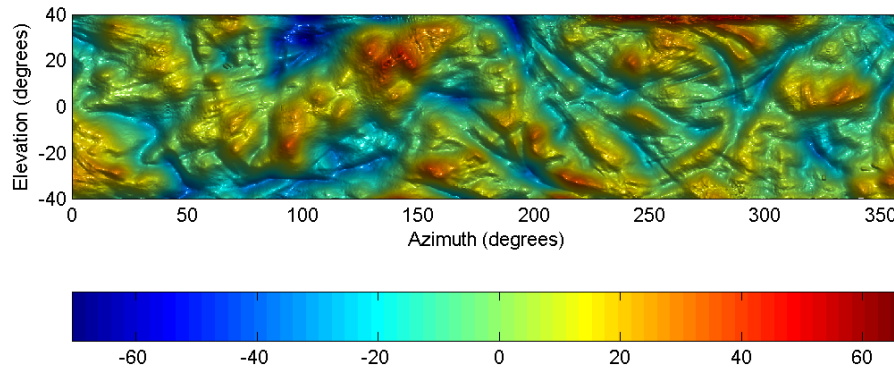


Figure 7.16: Measurements of the sphere, corrected for radial distortion, at ten degree increments in rotation viewed in a common co-ordinate system.

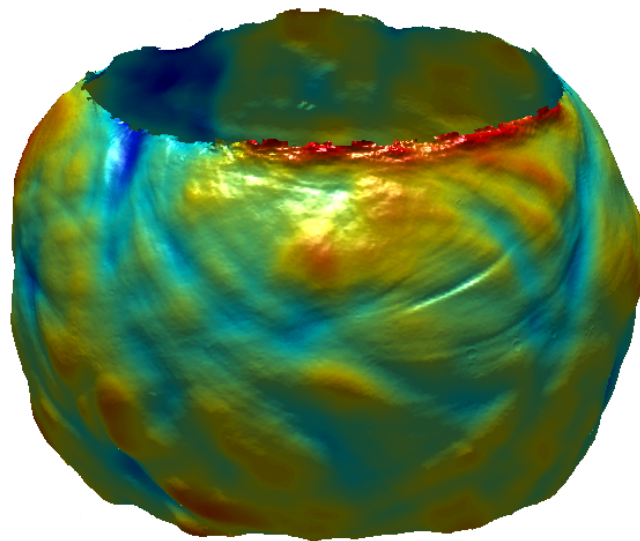


Figure 7.17: Departure from a perfect spherical surface plotted as a function of spatial position (but greatly exaggerated in scale).

7.3.3 Summary

This example has shown that visualisation can be used to verify the information on which data analysis is based. It is valuable to plot data in a variety of graphical forms because it provides a means for identifying erroneous assumptions about the data and mistakes in the implementation of algorithms used in the processing of the data.

7.4 Visualising the results of continuous modelling

7.4.1 Introduction

The data set for this visualisation arises from a problem in underwater acoustics [43]. An underwater acoustic transducer operating at a single frequency is used to generate a pressure distribution within a tank of water. The pressure distribution is measured on a set of coaxial open-ended cylinders of various radii and lengths. Specifically, measurements p_{ijk} of pressure are made on a set of concentric cylinders at points with cylindrical co-ordinates (r_i, θ_j, z_k) . For each value r_i of cylinder radius, the same values θ_j of angular co-ordinate are used for each value z_k of ‘height’ on the cylinder.

The pressure distribution created by a single-frequency acoustic source in a uniform medium is described by the Helmholtz equation. The pressure measurements are used to define boundary conditions on a cylinder and the Helmholtz equation in its surface integral form is solved to give a value of the pressure calculated at any point external to the cylinder. If values of pressure are calculated at the positions (r_i, θ_j, z_k) , then the experimental measurements and the calculated results can be compared directly.

The aim of the visualisation is to help the metrologist to compare the experimental measurements to the calculated results. In general, the metrologist is most interested in the location and relative sizes of the key features of the pressure distribution, so it is important that the visualisation enables those key features to be identified and compared. The Helmholtz equation treats pressure as a complex-valued quantity (for mathematical convenience) but in general the metrologist is interested in the magnitude of the pressure rather than the phase. Hence, the quantity to be visualised is a single-valued function of three variables (the cylindrical polar co-ordinates). In general, the metrologist is interested in relative values rather than absolute values, so it can be useful to normalise the magnitude of pressure before the quantity is visualised.

Some acoustic transducers are very directional, meaning that the pressure is concentrated in one region of space. The location of this concentration, the main lobe or main beam, is an important feature. Many transducers have side lobes as well, which are peaks away from the main lobe, often caused by interference patterns. The number, location, and size relative to the main lobe of the side lobes are all important features.

7.4.2 Visualisations

The first set of visualisations are shown in Figure 7.18. These visualisations show the measured and calculated pressure magnitudes as coloured contour plots in three-dimensional space. They enable the user to compare qualitatively the two sets of data without closely examining the details, particularly since the two plots are not quite to the same colour scale. However, from these plots the user can see that the two main lobes and the side lobes predicted by the calculations are in approximately the same place as those that are measured.

Figure 7.20 shows the differences between the two plots in Figure 7.19 as a flattened contour plot. The darkest areas are the areas where the differences are largest. It is clear

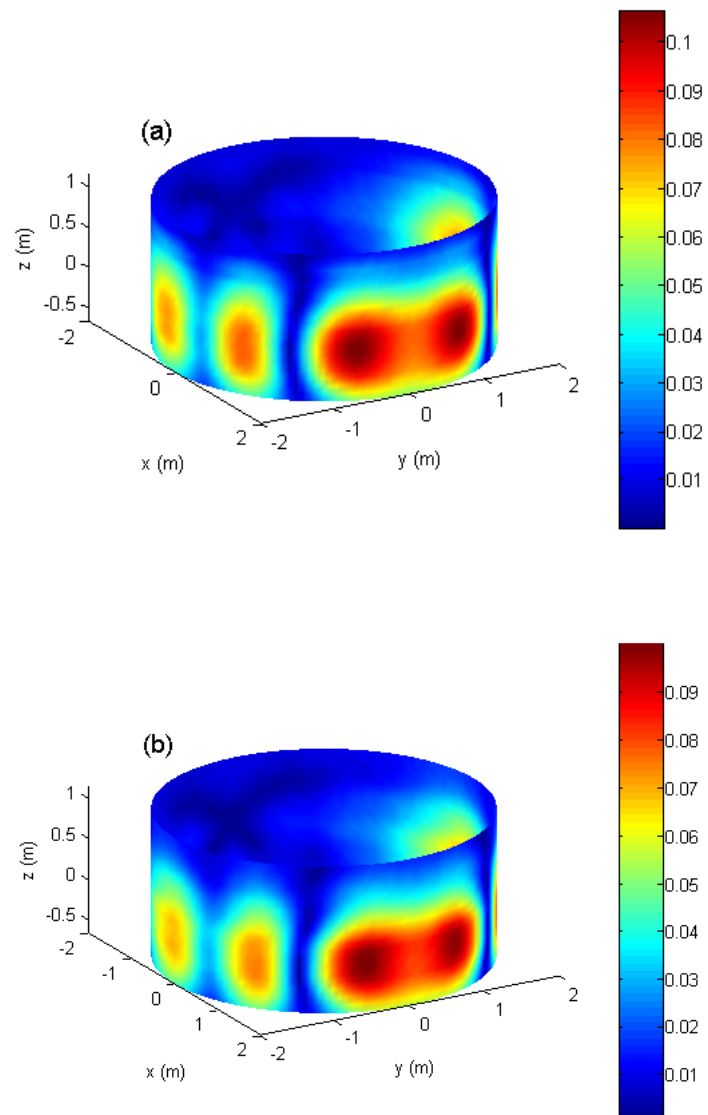


Figure 7.18: Colour contour plots of (a) measured and (b) calculated pressure magnitude.

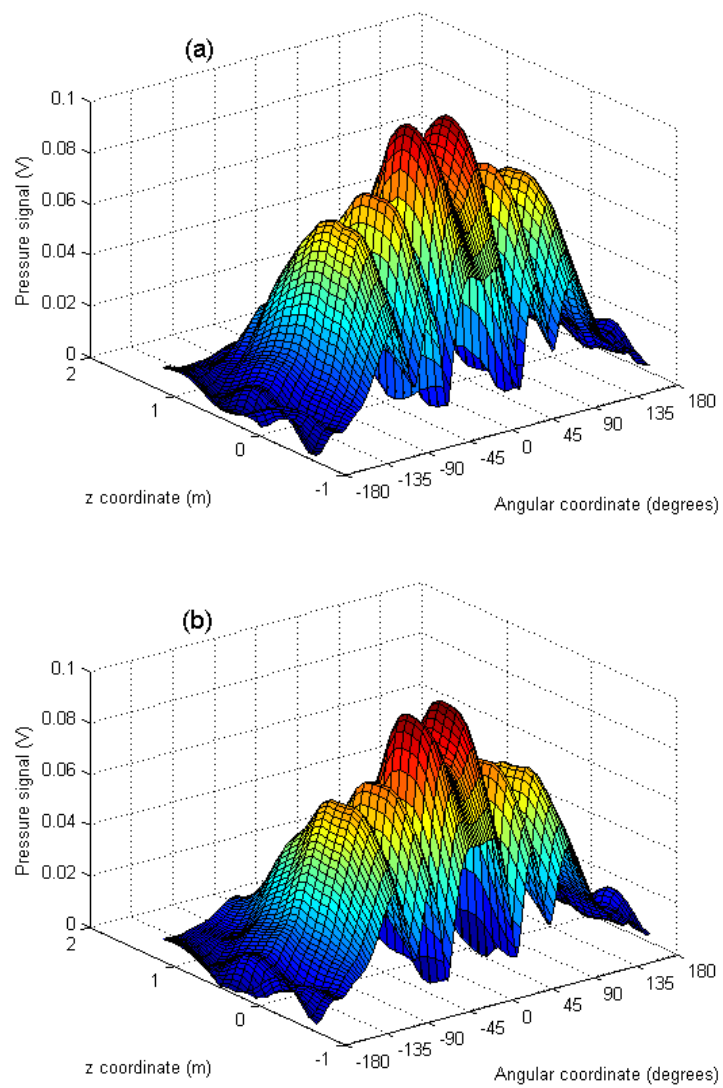


Figure 7.19: Surface plots of (a) measured and (b) calculated pressure magnitude.

that the differences are largest in one particular angular region for all values of the vertical coordinate. Efforts to improve the accuracy of the predictions would focus on the cause of this property of the differences.

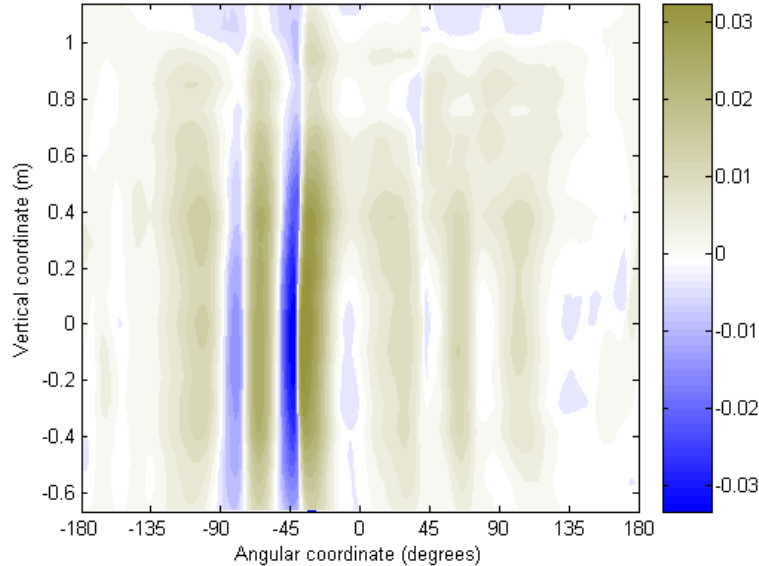


Figure 7.20: Differences between the plots shown in Figure 7.19.

The visualisations in Figure 7.18 are limited in scope, because the user cannot see all the data simultaneously. Since the values shown lie on a cylinder of constant radius, only z and θ are varying, so the data can be plotted as a surface plot as shown in Figure 7.19. These plots enable the user to see more of the data than is possible using the plots in Figure 7.18.

The plots in Figure 7.19 indicate that the number, size and position of the main and side lobes are about the same in the two data sets. The plots also make the differences between the measured and calculated data sets clearer than the three-dimensional plots in Figure 7.18. In particular, it is clear from the height of the surface relative to the grid lines that the calculated peak pressure on the main lobes is slightly lower than the measured value. However, these plots still do not make a quantitative visual comparison of the two data sets straightforward.

The most direct way of producing a quantitative comparison between the measured and calculated data sets is to create a line plot for each of the values z_k of z , plotting the measured and calculated values against the angular co-ordinates θ_j . A typical example is shown in Figure 7.21. This figure shows the measured and calculated values for $z_k = 0.19$ m. The agreement between measured and calculated values is quite good (the maximum difference is about 10 % of the measured value), with the lobes appearing in approximately the right places. A series of plots like this for different values z_k of z could be formed into an animation to show how the comparison varies with depth.

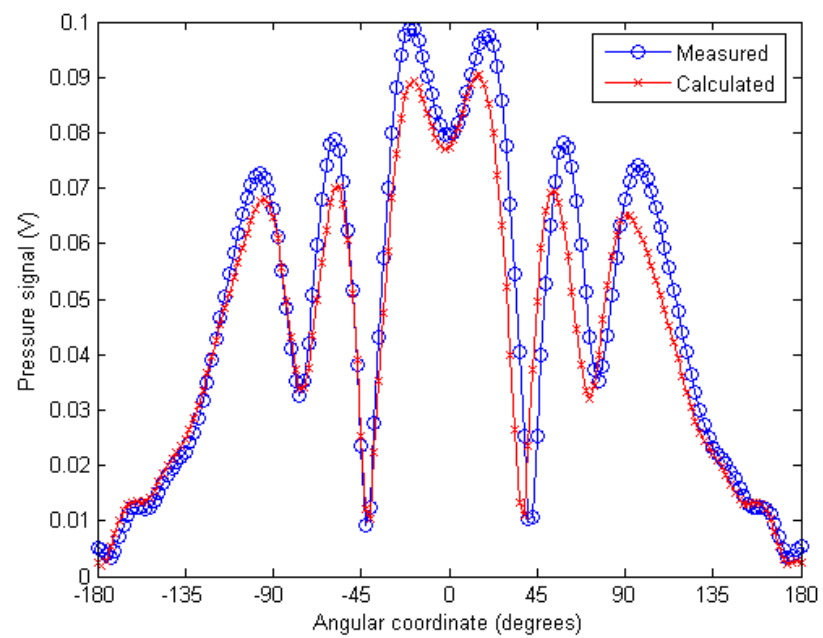


Figure 7.21: Line plot of a single circle of measured data and calculated results.

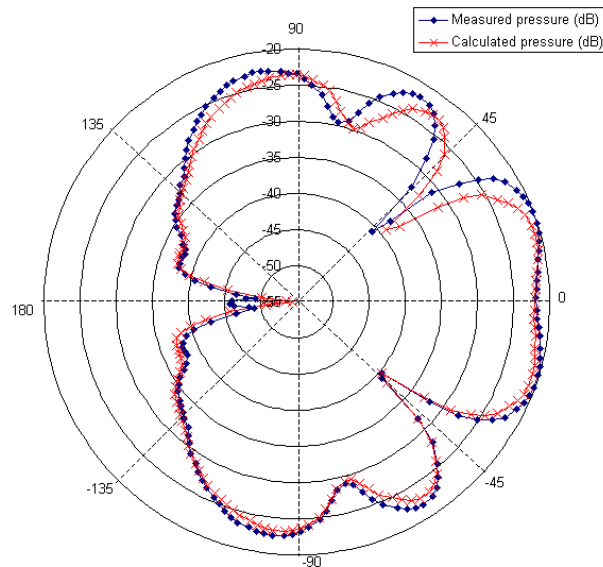


Figure 7.22: Polar plot of a single circle of measured data and calculated results. Amplitudes are plotted in decibels.

Figure 7.22 shows the comparison as a polar plot using decibel units (a logarithmic scale) for the magnitude. This type of plot is commonly used in acoustics applications, particularly for sonar and ultrasound, where the main lobe width strongly affects the performance of a device. The data are usually normalised, since only the relative sizes of the lobes are important. The polar plot gives a clear picture of the overall directional behaviour of the acoustic device, and the use of decibels emphasises the peaks of the pressure distribution. Experts in acoustics are accustomed to interpreting such plots, and it is important to visualise the data in a form with which the user is familiar.

7.4.3 Summary

This example has shown that different visualisations can provide different levels of information about the features of a data set that are of interest. It also shows that the best visualisation of a data set is often the one that the user is most familiar with and that identifies the features of interest to the user.

7.5 Visualising ordinal and categorical data

7.5.1 Introduction

As was mentioned in section 3.1, most metrology data is continuous. However, categorical data and ordinal data are still useful. Intercomparisons between different laboratories or different measurement devices include categorical data, and many cases of measurements

taken at varied intervals over a long time period are treated as ordinal data.

This example shows different ways of visualising ordinal and categorical data, along with some indication of ways to avoid pitfalls and produce better results. The aim of the visualisation is to present the data in a form that can be understood at a glance.

The data set used in the example was generated by a simple survey of the distribution of shoe sizes within a population of 77 people. In the initial study, only integral sizes were reported, with half-sizes being rounded down - e.g. $5\frac{1}{2}$ becomes 5. The initial results are shown in Table 7.1. In a second survey of the same group, half-sizes are allowed. The results of this survey are shown in Table 7.2, where it can be seen that, for example, the three people reporting a size of 5 in the first survey are now split into one size of 5 and two of $5\frac{1}{2}$.

Shoe Size	5	6	7	8	9	10	11	12
Frequency	3	5	14	13	19	17	5	1

Table 7.1: Results of the first survey of shoe sizes.

Shoe Size	5	6	7	8	9	10	11	12	$5\frac{1}{2}$	$7\frac{1}{2}$	$9\frac{1}{2}$
Frequency	1	5	6	13	11	17	5	1	2	8	8

Table 7.2: Results of the second survey of shoe sizes, including half sizes.

7.5.2 Visualisations

The visualisations shown in this section have been generated using Microsoft® Excel which, although best-known as a popular spreadsheet and calculation application, also includes the facility for producing charts and plots from the spreadsheet contents via its so-called chart wizard. This makes the visualisation of categorical data very easy: the values are entered into the cells of a spreadsheet, the columns and rows of interest are interactively highlighted and the wizard is used to produce the plot. Features of the chart such as axis labels, font sizes and data ranges can be edited interactively; in addition, the plot can be printed from within Excel or transferred to other packages such as Powerpoint or Word. Finally, because the plot is linked to the data in the spreadsheet, changes in the data are immediately reflected in the plot, which can be helpful when performing a “what-if” analysis.

If the data in the original survey, as in Table 7.1, are plotted as a bar chart, the results are as shown in Figure 7.23. If the same type of chart is created using the data from the second survey, exactly as shown in Table 7.2, then Figure 7.24 is produced. A problem is immediately apparent: the ordering of the data on the chart reflects the order in which the values appeared in the spreadsheet - i.e. the new entries for $5\frac{1}{2}$, $7\frac{1}{2}$ and $9\frac{1}{2}$ come after the previous entries. Whilst this might be desirable in some circumstances, it is confusing in this case because the data should be treated as ordinal rather than categorical.

Reordering the data so that the values are in order of increasing size and plotting as a bar chart produces Figure 7.25. Now a second problem, which is also present in Figure 7.24 can be seen: the spacing between the values on the Size axis does not reflect the difference between the values themselves. In fact, the physical spacing between successive values is

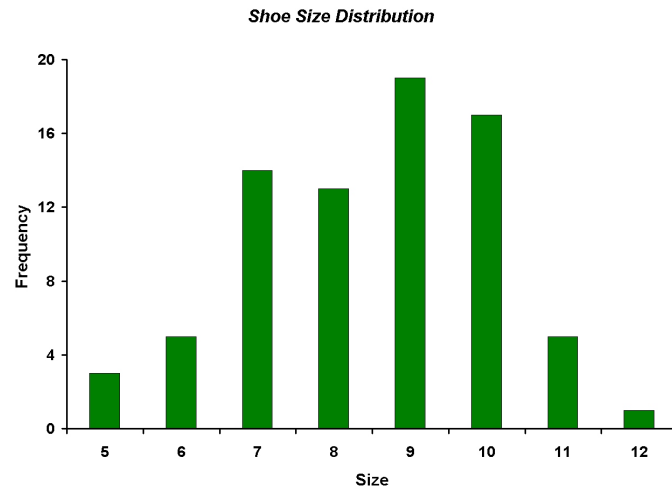


Figure 7.23: Using Microsoft® Excel to display the data in Table 7.1 as a bar chart.

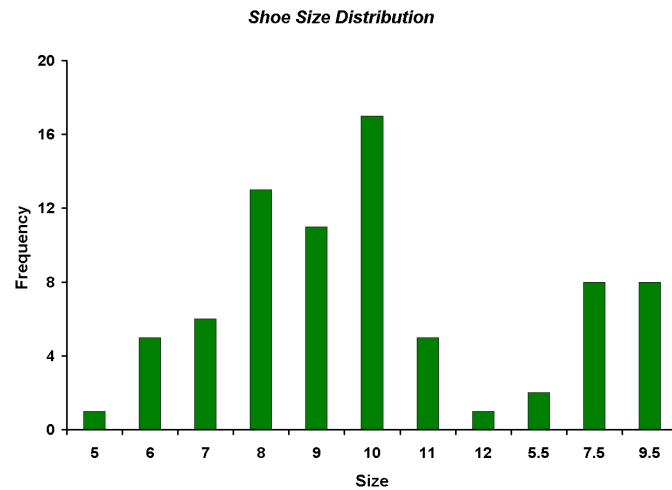


Figure 7.24: Displaying the data in Table 7.2 as a bar chart in Excel. Note that the ordering of the sizes reflects that of the dataset.

everywhere the same; for example, the first and second entries (whose difference in value is 0.5) are separated by the same distance as the third and fourth entries (whose difference is 1.0). Paying close attention to the labels on the axes would get around this problem, but this could be hard to do with more complicated datasets, as will be shown in section 7.6. In the present example, looking for size values with zero frequency (which could be missing data) is more difficult to do in a visualisation such as that shown in Figure 7.25, and it would perhaps be better to display datasets such as these without display artefacts, since they could mask interesting behaviour in the data itself.

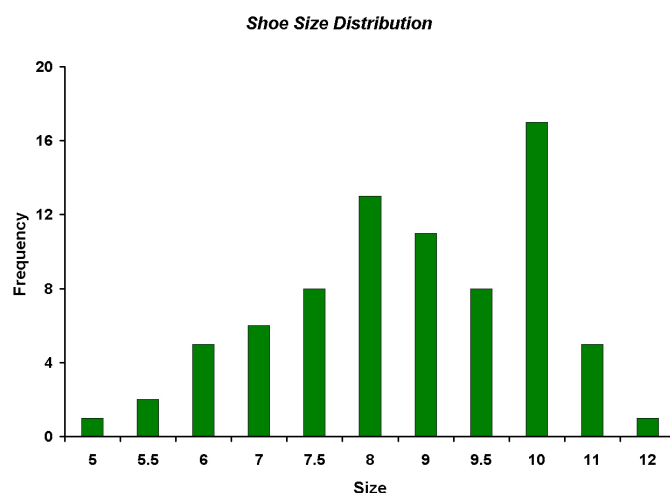


Figure 7.25: Displaying the data in Table 7.2, after sorting in order of increasing size, as a bar chart in Excel.

The problems of the ordering of values on the X axis and the physical separation between successive values on the X axis also affect its line chart and surface plot (see section 7.6). The problems arise because there are two different types of X axis in Excel charts, named Value and Category. Using a Value axis, the data is treated as continuously varying numerical values and the marker for a data point is placed at a location along the axis which reflects its value. Using a Category axis, the data is treated as a sequence of non-numerical text labels, the marker is placed at a location on the axis which reflects the point's position in the sequence, and the points are distributed evenly along the axis. These differences mean that a Category axis is only suitable for categorical data, or ordinal data that has already been sorted into the correct order, and if a data set does not fall into either of these types then a Value axis should be used.

The problems illustrated above arise because the bar chart (and most of Excel's built-in chart types) uses a Category X axis. By contrast, the XY (Scatter) chart uses a Value X axis. Using it to plot the contents of Table 7.2, we generate Figure 7.26. Now the data points are positioned correctly.

The problems with the treatment of the X axis in Excel's bar chart can be solved by using the NAG Schools Excel Add-in N-SEA. This is an Excel Add-in developed by NAG to support instruction in statistics and includes functionality for data sampling as well as the construction of frequency plots, histograms and box and whisker plots. Its plotting facility gives the option of ordering the data points according to their X or Y value, as well as

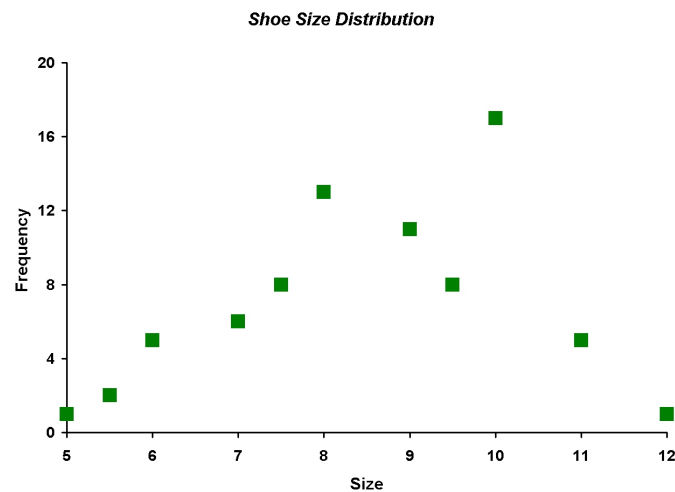


Figure 7.26: Displaying the data in Table 7.2 as a scatter plot in Excel. Note that the points are correctly positioned along the Size axis.

including data points with zero Y value in the plot. Using N-SEA to plot the data in Table 7.2 gives Figure 7.27; comparing this with Figure 7.25 shows this to give a clearer picture of the data. In particular, it is easy to spot the sizes whose frequency is zero.

Other displays are possible using N-SEA. For example, suppose that the discrete data set of Table 7.2 is transformed into a continuous one by grouping sizes together using unequal group widths, as shown in Table 7.3. N-SEA can display this as a continuous bar chart with variable widths (see Figure 7.28). Once again, this is not possible to do in Excel without introducing artefacts that could be misleading.

Shoe Size	5-6.5	6.5-7.25	7.25-8.5	8.5-9.25	9.25-9.75	9.75-10.5	10.5-12
Frequency	8	14	13	11	8	17	6

Table 7.3: Data in Table 7.2 presented as a continuous variable

7.5.3 Summary

Some visualisation methods, such as bar charts, are intended for use with categorical data, and may produce misleading results when used with ordinal data. In particular, the ordering of data points and their spacing along the horizontal axis can strongly affect the user's interpretation of visualisations of ordinal data. Preparation of the data prior to visualisation can improve the subsequent images significantly, but some visualisation tools automatically carry out the necessary operations before creating an image.

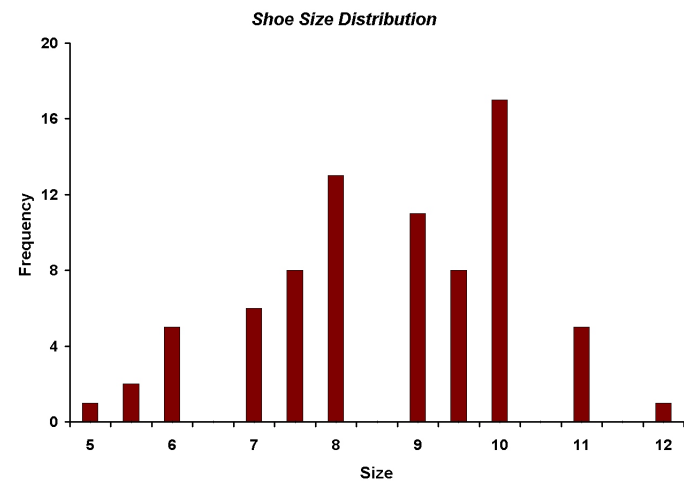


Figure 7.27: Displaying the data in Table 7.2 using N-SEA in Excel. The data is automatically ordered in ascending Size value by N-SEA, and the spacing between the points on that axis reflects their numerical value.

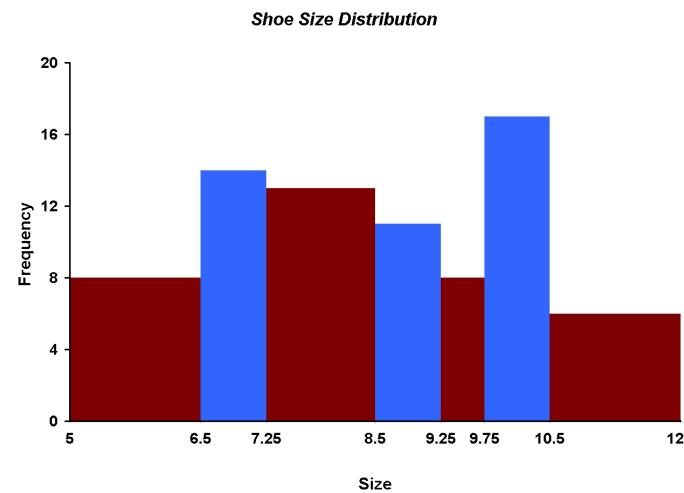


Figure 7.28: Using N-SEA in Excel to display the results of the second survey after it has been grouped into bins of varying width as in Table 7.3.

7.6 Validating a visualisation

7.6.1 Introduction

The first data set used in this example comes from the measurement [38] of fluid density inside a pore as a function of fluid pressure, called an adsorption isotherm. A bulk fluid undergoes a phase change from gas to liquid at its saturation point but, as the pressure is increased towards its value at saturation, a confined fluid may condense to a liquid phase before this point is reached. This is the phenomenon of capillary condensation, and it shows up in the adsorption isotherm as a vertical jump in density at the pressure where the two phases (gas and liquid) coexist in the pore. The aim of the visualisation is to identify the pressure at which capillary condensation occurs.

The second data set comes [35] from financial modelling, where calculations in the theory of options on swap agreements (so-called “swaptions”) using a model like Black-Scholes (a partial differential equation model of option values) produce values for the option volatility (the standard deviation of the change in value of the option) as a function of the underlying maturity of the swap and the expiry time on the option. The aim of the visualisation is to produce a clear display of volatility as a function of maturity and expiry.

7.6.2 Visualisations

Initially, the adsorption isotherm data set was plotted using an Excel line plot, as shown in Figure 7.29. The plot was then validated by creating a scatter plot of the same data. This plot is shown in Figure 7.30. Figure 7.30 clearly shows the irregular spacing between pressure values, arising from the experimentalist’s attempts to determine the transition point, along with the vertical jump, which is caused by the data set containing two density values at the same pressure. By contrast, Figure 7.29 produces a qualitatively different, and potentially misleading, result. The fixed spacing between successive pressure data does not reflect their value, and the vertical jump is missing in this representation of the data because the axis gives a misleading sense of gradient. If the viewer attempted to identify the transition point by interpolation from a plot similar to Figure 7.29, the result would be wrong.

The difference between the two plots is caused by the use of a Category axis in Figure 7.29 and a Value axis in Figure 7.30, as is explained in section 7.5.2. Category axes cannot display discontinuities in data such as the jump caused by the capillary condensation.

The problems associated with the Category axis in Excel are also present for more complex plots. The second data set can be plotted in Excel using the Surface plot, as shown in Figure 7.31. The visualisation is validated by displaying the same data set using IRIS Explorer [37], with the results shown in Figure 7.32. The labels on the Maturity and Expiry axes in Figure 7.31 show the same problems with the Category axis that were encountered in the line chart (Figure 7.29). The uneven spacing on the axis can lead to misinterpretation of the figure. For example in Figure 7.31, it is not clear whether the somewhat abrupt change in slope of the front edge of the surface at the 10 year maturity mark is a feature of the data set, or if it is caused by the distortion due to the five-fold change in spacing which takes place here. In Figure 7.32, the irregular spacing between

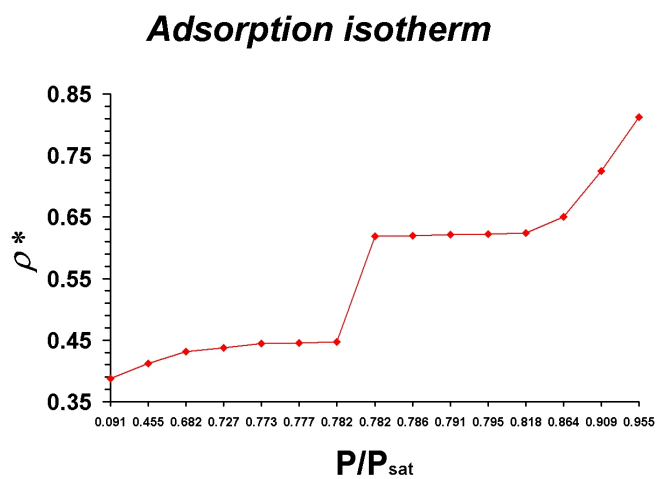


Figure 7.29: Displaying an adsorption isotherm for fluid in a pore using a line plot in Excel.

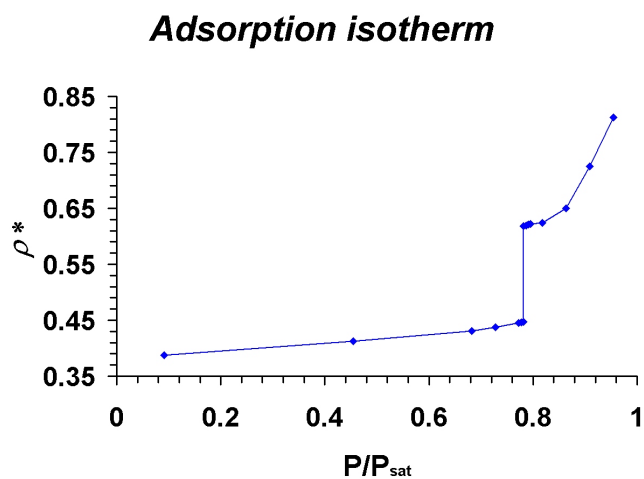


Figure 7.30: Plotting the same data as shown in Figure 7.29 as a scatter plot.

successive values for the independent variables is correctly reflected in the visualisation, and it is possible to get a better idea of the shape of the volatility surface [36].

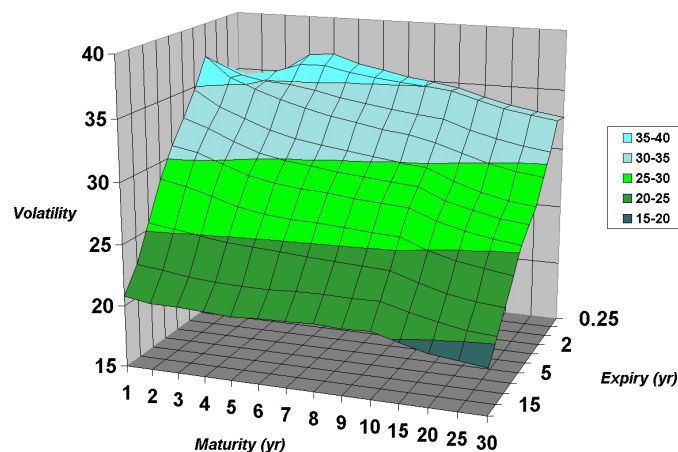


Figure 7.31: Volatility surface data visualised using a surface plot in Excel.

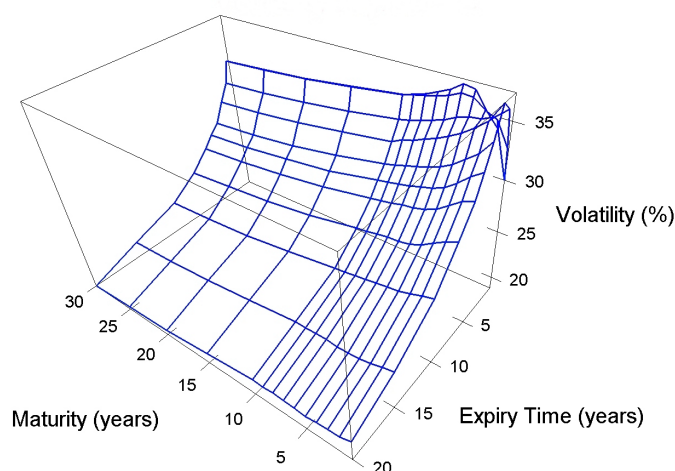


Figure 7.32: The same data as shown in Figure 7.31 using IRIS Explorer.

The use of Category axes by the Excel surface plot means that it is unable to display discontinuities in the surface (the two-dimensional analogue of the vertical jump in the plot in 7.30). This is important in the present context because discontinuities can occur when using the Black-Scholes equation to model the value of an option as a function of stock price and time to maturity. At certain maturities, a discrete dividend is paid, causing the value of the option to fall. This fall should show up as a vertical step in the surface - but it does not appear when the data is plotted in Excel, which makes no allowance for changes in the value of the spacing between independent variables. For a more correct representation of the volatility surface, other visualisation packages such as IRIS Explorer must be used.

7.6.3 Summary

Validating visualisations by comparison with the results of an alternative visualisation method or an alternative software package can highlight deficiencies in the original images. Validation can also suggest methods for improving subsequent visualisations of similar data sets.

Chapter 8

Summary

This Guide provides information on good practice in the use of visualisation techniques applied to metrology data. It has been prepared to help metrologists and others who are concerned with the generation and interpretation of numerical data in the effective use of visualisation techniques. In particular, the aims of the Guide are

- to show the benefits of visualisation,
- to give practical advice on how to perform visualisation, and
- to make existing good practice available in a readily accessible form.

Visualisation is about communicating complex information graphically. The concern here is with information that takes the form of numerical data, and visualisation is undertaken to explore the data and to present the data, or information derived from it, to others in an accessible form.

For visualisation to be successful, it is important that the objectives of, and constraints on, the visualisation are carefully considered. Delivering a visualisation of the wrong data, or one showing the wrong characteristics or attributes of the data, or one which demands inappropriate time or resources to develop or use, is worthless. Common mistakes made are not being clear why a visualisation is being used and not being able to judge whether a visualisation has been successful.

Being able to answer the following questions can help to clarify the objectives of, and constraints on, the visualisation:

- What is the purpose of the visualisation?
- What form will the visualisation take?
- Who is the audience for the visualisation?
- What is the nature of the data?
- What visualisation software is available?

- Does the visualisation need to be validated?

In considering the data to be visualised, it is important to consider

- How many dimensions are there within the data?
- What structure is inherent in the data?
- What part of that structure should be preserved by the visualisation?

The basis for visualisation is to find an effective mapping from the data to appropriate visualisation dimensions: these are characteristics of objects such as type, location, angle, size, colour hue, and colour saturation. The ability to perceive the positions of objects and their physical relationships is fundamental to human perception, and this makes the spatial dimensions of a visualisation the primary way of conveying information. Irrespective of the number of dimensions within the data, data are almost always visualised in two dimensions because computer screens and printed media are usually two-dimensional. Perceptual cues, including occlusion, lighting and assumptions about orthogonality, are often used to infer depth in a visualisation that is reliant on spatial dimensions.

Other examples of visualisation dimension include colour and time. The basic perceptual properties of colour are hue, saturation and value. Different colour maps are obtained by encoding information differently in terms of these properties. An appropriate choice of colour map can greatly improve the effectiveness of a visualisation using colour. Animations, including movies, arise from the mapping of a dimension within the data to time. Visualisations using time can exploit the sensitivity of the human visual system to movement, but an animation that lasts a long time may require the observer to remember parts of the visualisation.

Combining visualisation dimensions, for example spatial dimensions with colour, can help to emphasise and enhance the information conveyed by a visualisation. Another approach is to provide interaction between the visualisation and the observer. There are different types of interaction: allowing the observer to control how the scene is viewed, for example by allowing camera panning, and allowing the observer to modify parameters of the data that is displayed.

It is essential that the software that produces visualisations be validated and tested, and that the results of the validation and testing activities be documented and made available to the users of visualisation software. The validation of visualisations needs to address the issues of

- the manner in which the visualisation software represents and transforms the data, and
- the way in which the user of the software understands and interprets the results of the visualisation.

The development of visualisation software should thus be approached in the same manner as any other software development project, with the preparation of user and functional specifications and the definition of a comprehensive test plan to address the above issues.

Uncertainty evaluation is a topic of particular concern to metrologists. In the context of uncertainty evaluation, the objectives of visualisation include communicating and comparing information, expressed probabilistically, about the values of (measured) quantities. The objects to be visualised include probability distributions, and information derived from a probability distribution, including a best estimate, an associated standard uncertainty, a coverage interval, and samples drawn randomly from the distribution. Many of the visualisation dimensions discussed above are useful for this purpose.

Much continuous modelling software, such as that used for finite element and finite difference analysis and computational fluid dynamics, incorporates visualisation for displaying both a model and the results obtained from the model. Visualisation will appear to many users of such software to offer a succinct and cost-effective way of studying the results obtained from a model. However, it is important for the user to consider the visualisations produced by such software critically: the visualisations can be affected by interpolation, extrapolation, smoothing, choice of co-ordinate system, choice of scaling factors and display tolerances, that are undertaken by the software and may not be transparent to the user.

Finally, the Guide presents a number of examples intended to illustrate the visualisation techniques discussed. The examples consider the visualisation of experimental data, uncertainties associated with measured data, results of a continuous model, ordinal and categorical data, and the validation of a visualisation.

Bibliography

- [1] R. M. Barker, M. G. Cox, A. B. Forbes, and P. M. Harris. Best Practice Guide No. 4. Discrete modelling and experimental data analysis. Technical report, National Physical Laboratory, Teddington, UK, 2004. Available for download from the SSfM website at www.npl.co.uk/ssfm/download/bpg.html. 54
- [2] J. Barrett, M. G. Cox, M. P. Dainton, and P. M. Harris. A methodology for testing the numerical accuracy of scientific software used in metrology. In P. Ciarlini, M. G. Cox, E. Filipe, F. Pavese, and D. Richter, editors, *Advanced Mathematical Tools in Metrology V. Series on Advances in Mathematics for Applied Sciences Vol. 57*, pages 31–40, Singapore, 2001. World Scientific. 54
- [3] R. W. Beatty. Insertion loss concepts. *Proc. IEEE*, 6:663–671, 1964. 91
- [4] BIPM, IEC, IFCC, ILAC, ISO, IUPAC, IUPAP, and OIML. Guide to the Expression of Uncertainty in Measurement, Supplement 1, Numerical Methods for the Propagation of Distributions. To be published. 57, 60, 65
- [5] BIPM, IEC, IFCC, ISO, IUPAC, IUPAP, and OIML. Guide to the Expression of Uncertainty in Measurement, 1995. ISBN 92-67-10188-9, Second Edition. 57, 96
- [6] M. Botts, S. Uselton, J. Walton, H. Watkins, D. Watson, and H. Rushmeier. Metrics and benchmarks for visualization. In *Visualization '95*, pages 622–626. IEEE Comput. Soc. Press, 1995. 51
- [7] R. Brooks, J. Wilson, and R. Thorpe. Geometry unifies process control, production control and alarm management. *Computing & Control Engineering Magazine*, February 2004. 40
- [8] M. G. Cox. The evaluation of key comparison data. *Metrologia*, 39:589–595, 2002. 58
- [9] M. G. Cox, T. J. Esward, P. M. Harris, N. McCormick, N. M. Ridler, M. J. Salter, and J. P. R. B. Walton. Visualisation of uncertainty associated with classes of electrical measurement. Technical Report CMSC 44/04, National Physical Laboratory, Teddington, UK, 2004. 3, 59, 60, 92, 95
- [10] M. G. Cox and P. M. Harris. Best Practice Guide No. 6. Uncertainty evaluation. Technical report, National Physical Laboratory, Teddington, UK, 2004. Available for download from the SSfM website at www.npl.co.uk/ssfm/download/bpg.html. 57, 59
- [11] M. G. Cox and P. M. Harris. Software specifications for uncertainty evaluation. Technical Report CMSC 40/04, National Physical Laboratory, Teddington, UK, 2004. 60

- [12] C. F. Dietrich. *Uncertainty, Calibration and Probability*. Adam Hilger, Bristol, UK, 1991. 57
- [13] M.J. Downs, K.P. Birch, M.G. Cox, and J.W. Nunn. Verification of a polarization-insensitive optical interferometer system with subnanometric capability. *Precision Engineering*, 17(1), 1995. 82
- [14] T. Esward, G. Lord, and L. Wright. Model validation in continuous modelling. Technical Report CMSC 29/03, National Physical Laboratory, Teddington, UK, 2003. 54, 74
- [15] L. Fox and I.B. Parker. *Chebyshev Polynomials in Numerical Analysis*. Oxford University Press, 1968. 83
- [16] J. Gilby and J. Walton. Good Practice Guide No. 13. Data visualization. Technical report, National Physical Laboratory, Teddington, UK, 2003. Available for download from the SSfM website at www.npl.co.uk/ssfm/download/bpg.html. 3
- [17] A. Globus and S. Uselton. Evaluation of visualisation software. *Comput. Graph.*, 29:41–4, 1995. 51, 54
- [18] A. Inselberg, B. Dimsdale, A. Chatterjee, and C.-K. Hung. Parallel co-ordinates: survey of recent results. In J. P. Allebach and B. E. Rogowitz, editors, *Proceedings of SPIE - Volume 1913 Human Vision, Visual Processing and Digital Display IV*, pages 582–599, 1993. 40
- [19] D. M. Kerns and R. W. Beatty. *Basic Theory of Waveguide Junctions and Introductory Microwave Network Analysis*. Pergamon Press, London, 1967. 91
- [20] M.J. Korczynski, M.G. Cox, and P.M. Harris. Convolution and uncertainty evaluation. In *Advanced Mathematical and Computational Tools in Metrology VII*. World Scientific, 2005. To appear. 57
- [21] A. Lopes and K.W. Brodlie. Accuracy in 3D particle tracing. In H.-C. Hege and K. Polthier, editors, *Mathematical Visualization: Algorithms, Applications and Numerics*, pages 329–341. Springer Verlag, 1998. <http://www.scs.leeds.ac.uk/vis/adriano/acc-particle.ps.gz>. 54
- [22] Sira Ltd. Visual modelling and data visualization. Technical report, Sira Ltd., 1999. Available for download from the SSfM website at www.npl.co.uk/ssfm/download/documents/mpu_8-53-1.pdf. 3
- [23] A. T. Pang, C. M. Wittenbrink, and S. K. Lodha. Approaches to uncertainty visualization. *The Visual Computer*, 13:370–390, 1997. 52, 55
- [24] A. Papoulis. *Probability, Random Variables and Stochastic Processes*. McGraw-Hill, New York, 1984. 60
- [25] N. M. Ridler. A review of existing national measurement standards for RF and microwave impedance parameters in the UK. *IEE Colloquium Digest*, 99/008:6/1–6/6, 1999. 96
- [26] N. M. Ridler and M. J. Salter. Propagating S-parameter uncertainties to other measurement quantities. In *58th ARFTG (Automatic RF Techniques Group) Conference Digest*, 2001. 91

- [27] N. M. Ridler and M. J. Salter. An approach to the treatment of uncertainty in complex S-parameter measurements. *Metrologia*, 39:295–302, 2002. 91
- [28] E. R. Tufte. *Envisioning Information*. Graphics Press, Connecticut, 1990. 54
- [29] E. R. Tufte. *The Visual Display of Quantitative Information, 2nd edition*. Graphics Press, Connecticut, 2001. 54
- [30] J. W. Tukey. *Exploratory Data Analysis*. Addison-Wesley, Reading, MA, 1977. 54
- [31] S. Usselton, G. Dorn, C. Farhat, M. Vannier, K. Esbensen, and A. Globus. Validation, verification and evaluation. In *Visualization '94*, pages 414–418, 1994. 51
- [32] J. P. R. B. Walton. Visual benchmarking: A practical application of 3D publishing. In H. Jones, R. Raby, and D. Vicars, editors, *Proceedings of Eurographics UK 14th Annual Conference, Vol. 2*, pages 339–351, 1996. <http://www.nag.co.uk/doc/TechRep/PS/tr9.96.ps>. 54, 72
- [33] J. P. R. B. Walton and M. Dewar. See what I mean? Using graphics toolkits to visualise numerical data. In H.-C. Hege and K. Polthier, editors, *Visualization and Mathematics: Experiments, Simulations and Environments*, pages 279–299. Springer Verlag, 1997. <http://www.nag.co.uk/doc/TechRep/PS/tr8.96.ps>. 54
- [34] J.P.R.B. Walton. Visual benchmarking: A practical application of 3D publishing. In H. Jones, R. Raby, and D. Vicars, editors, *Proceedings of Eurographics UK 14th Annual Conference, Vol. 2*, pages 339–351, 1996. Available for download from www.nag.co.uk/doc/TechRep/PS/tr9.96.ps. 52
- [35] J.P.R.B. Walton. Get the picture: visualizing financial data - part 1. *Financial Engineering News*, 36, 2004. 113
- [36] J.P.R.B. Walton. Get the picture: visualizing financial data - part 2. *Financial Engineering News*, 37, 2004. 115
- [37] J.P.R.B. Walton. NAG's IRIS Explorer. In C.D. Hansen and C.R. Johnson, editors, *The Visualization Handbook*, pages 615–636. Academic Press, 2005. Available for download from www.nag.com/IndustryArticles/ch32.pdf. 113
- [38] J.P.R.B. Walton and N. Quirke. Capillary condensation: a molecular simulation study. *Molecular Simulation*, 2:361–391, 1989. 113
- [39] K. Weise and W. Wöger. A Bayesian theory of measurement uncertainty. *Measurement Sci. Technol.*, 3:1–11, 1992. 59
- [40] W. Wöger. Probability assignment to systematic deviations by the Principle of Maximum Entropy. *IEEE Trans. Instr. Measurement*, IM-36:655–658, 1987. 59
- [41] L. Wright. Visualisation, uncertainties and continuous modelling. Technical Report CMSC 39/04, National Physical Laboratory, Teddington, UK, 2004. 3, 71, 72, 73
- [42] Louise Wright. An introduction to the parallel coordinate technique applied to measurement data. Technical Report ???, National Physical Laboratory, Teddington, UK, 2006. 42, 43

- [43] Louise Wright, Stephen Robinson, Victor Humphrey, Peter Harris, and Gary Hayman. The application of boundary element methods to near-field acoustic measurements on cylindrical surfaces at NPL. Technical Report DQL-AC 011, National Physical Laboratory, Teddington, UK, 2005. 102

Appendix A

Glossary

Animation A sequence of images displayed one after another. The difference between successive images is usually related to either (a) a change in one of the independent variables in the visualisation, or (b) a new camera position within the three-dimensional scene. (a) is one way of showing how the visualisation depends on another independent variable (which could be time); (b) gives more information about a three-dimensional scene by showing it from different viewpoints.

Camera dollying Moving the camera closer to, or further away from, the scene.

Camera panning Moving the camera in a plane normal to the line connecting its original location to the centre of the scene.

Camera rotation Moving the camera so that it always points at the centre of the scene and the distance between it and the scene is held constant.

Categorical data Discrete, unordered data – e.g., Bill’s results, Jeff’s results.

Colour hue The component of colour related to its dominant wavelength - e.g. red, blue.

Colour map A function relating data value to colour. A common visualisation technique for scalar data is to map values to colour and display the colours, usually alongside a key or legend which shows the correspondence between colour and value.

Colour saturation The purity – i.e., the amount of hue present – in a colour. Complete saturation corresponds to a pure hue; zero saturation is greyscale (i.e., white, black, and all the greys in between).

Colour value The brightness of the colour. Dark colours have low values; bright colours have high values.

Continuous data Values which are not selected from a discrete set, typically expressed as floating point numbers – e.g., temperature.

Contour A geometric construction that joins all points in a scene where a data variable has the same value. For data sets with two independent variables (coordinates), this is a line, while for data sets with three coordinates, this is a surface (also known as an isosurface).

Dimension of data set The number of independent variables in the data set. The dimension of the data set usually determines the type of visualisation technique that can be applied to the data set.

Dimension of visualisation See **Visualisation dimension**.

Explicit model A relation between variables that gives one (e.g., z) as a function of the others (e.g., x and y). The latter are then independent variables, while the former is the dependent variable.

Focus When a person looks at a point, the muscles in the eye alter the focus of the eye to obtain the sharpest image possible. The sensing of the change in focus provides a perceptual cue to the distance of the point being observed.

Glyph A discrete geometric object whose properties are mapped from data values. For example, an arrow can be used to show the orientation and strength of a vector field at a point.

Implicit model A relationship between variables that equates some function of them to a constant.

Interactive visualisation Being able to change the viewing parameters, and/or interrogate elements in the scene in order to gain increased insight into the data being visualised.

Isosurface See **Contour**.

Occlusion The concealing of elements in the scene by other elements or objects that are closer to the viewpoint.

Ordinal data Discrete, ordered data – e.g., Monday’s results, Tuesday’s results.

Orthographic projection A type of parallel projection where the direction of projection is the same as the surface normal to the projection plane. Object size does not depend on distance from viewpoint; parallel lines remain parallel.

Parallax The apparent displacement of an object when seen from two different points (not on a straight line with the object). Can provide a perceptual clue to its distance from the viewpoint.

Perceptual cue An element in the visualisation which gives the viewer information about things like depth, positioning and relative size of objects in the scene. Examples include occlusion and lighting.

Perspective projection The complete projection model of a scene onto an image plane via a pinhole camera model. Object size depends on distance from viewpoint; parallel lines meet at a vanishing point.

Scalar quantity A quantity that has magnitude but no direction – e.g., temperature.

Scatter plot Visualisation technique used to display and compare n sets of data values by displaying them as points in n -dimensional space.

Scene An arrangement of graphical objects (usually) in three-dimensional space.

Slice Reducing the number of visualisation dimensions by holding one of the coordinates at a fixed value.

Specular reflection A sharply defined beam of light resulting from reflection off a smooth surface.

Stereopsis The process leading to visual perception of the depth or distance of objects, due to each eye seeing different images of the scene.

Tensor quantity A quantity defined relative to a set of coordinate axes that obeys certain rules when the coordinate axes are transformed. Tensor is commonly used to mean a tensor of rank two, such as stress or strain. Such quantities can be written as matrices. A vector is a tensor of rank one.

Vector quantity A quantity having direction as well as magnitude – e.g., velocity.

Vergence When a person looks at a point, the eyes are rotated so that the image of that point lies at the centre of vision of each eye. As the eyes are separated, this means that the relative rotation of the eyes varies with the distance to the imaged point. The sensing of relative eye rotation provides a perceptual cue to the distance of the point being observed. This cue is present when looking at real three-dimensional objects but not when observing two-dimensional representations.

View A two-dimensional image of a three-dimensional scene.

Visualisation dimension A feature of the visualisation that represents a property of the data. Examples include spatial dimensions, colour hue and saturation, and time.

Volumetric rendering The visualisation of three-dimensional volume data as a two-dimensional data set on an image plane. Each point of the two-dimensional data set is a function of the values of the volume data along a line through that point. Common functions used are the maximum value or mean value of the volume data along the line.

Appendix B

Multimedia objects

This document contains a number of multimedia objects for visualization. They are available by viewing the document electronically using *Acrobat Reader*, available from

<http://www.adobe.com/products/acrobat/readstep2.html>

To ensure the objects are visible, appropriate viewers are required to be installed.

Two types of multimedia object are used:

Movies provided with the document in the form of AVI movie files (section 5.2.4). The movies can be viewed using a suitable player such as *Windows Media Player*. The movies were created using the Indeo5 codec, which can be obtained from

<http://www.ligos.com/indeo.htm>.

Views of (three-dimensional) objects that allow also some interaction with the objects displayed (section 5.2.7). To interact with these views a web browser, such as Internet Explorer, is used to download a page containing a Java applet from the Internet or to open a page¹ provided with the document. To display the view a browser plug-in is required, such as *Cortona* or *Cosmoplayer*, dependent on the computing platform. *Cortona* can be obtained from

<http://www.parallelgraphics.com/products/cortona>

and is available for Windows and Macintosh operating systems. *Cosmoplayer* can be used in Unix environments, and can be obtained from

<http://www.sgi.com/software/cosmo/>

The web pages have been tested using Internet Explorer with the *Cortona* VRML viewer and the Microsoft Java VM. They have also been tested using Netscape 4.x, with the *Cosmoplayer* VRML viewer and the Netscape Java VM in an IRIX environment.

¹These pages are provided as files with extension '.wrl'.

Links to the multimedia objects provided in this document are given below. The objects may also be accessed from the sections indicated, where information on how to use the objects, and interpret what they show, is also provided.

- Movie 1 for co-ordinate system transformation (section 5.2.4)
- Movie 2 for co-ordinate system transformation (section 5.2.4)
- Movie 3 for co-ordinate system transformation (section 5.2.4)
- Movie 4 for co-ordinate system transformation (section 5.2.4)
- Web page 1 for co-ordinate system transformation (section 5.2.3)
- Web page 2 for co-ordinate system transformation (section 5.2.3)
- Web page 3 for co-ordinate system transformation (section 5.2.7)
- Web page for comparison loss in microwave power meter calibration (section 7.2).

Appendix C

Further reading and useful resources

Listed here briefly are some useful textbooks on the visualisation of data and some links to Internet resources on visualisation that readers of this Guide might find helpful.

C.1 Textbooks

The author whose works are regarded by many people as the key texts in this field is Edward Tufte. Information about his work and the courses he runs on visualisation can be found on his web site at <http://www.edwardtufte.com/tufte/>. His three best known textbooks are:

- The Visual Display of Quantitative Information, E R Tufte, Graphics Press, Cheshire, Connecticut, 1983
- Envisioning Information, E R Tufte, Graphics Press, Cheshire, Connecticut, 1990
- Visual Explanations: Images and Quantities, Evidence and Narrative, E R Tufte, Graphics Press, Cheshire, Connecticut, 1997

All the above books are worth studying for practical advice on communicating data visually and provide advice that can be used both in preparing graphs and charts for printed material and for computer display. A further pamphlet of his that may be of interest to readers who frequently have to present data and images using PowerPoint is:

- The Cognitive Style of PowerPoint, E R Tufte, Graphics Press, Cheshire, Connecticut, 2003

Other useful texts are:

- Visualizing Data, W S Cleveland, Hobart Press, Summit, New Jersey, 1993

This text is full of practical advice on the production and use of graphs for understanding data with examples of good and bad practice. It is a good companion to the books by Tufte. In particular, it makes extensive use of multiway dot plots.

- The Visualization Handbook, C.D. Hansen and C.R. Johnson (eds), Academic Press, 2005

This book provides an overview of the field of visualisation by presenting its basic concepts, providing a snapshot of current visualisation software systems and examining research topics that are advancing the field. It is intended for a broad audience, including not only the visualisation expert seeking advanced methods to solve a particular problem, but also the novice looking for general background information on visualisation topics.

- Plain Figures, Myra Chapman and Basil Mahon, Cabinet Office (Management and Personnel Office), Civil Service College, HMSO, London, 1986

This book provides short, clear advice on presenting data graphically and despite its publication date is still useful for understanding how to communicate effectively with graphs, especially for social and economic statistics.

- Visualization: Using Computer Graphics to Explore Data and Present Information, Judith R Brown, Rae Earnshaw, Mikael Jern and John Vince, John Wiley, 1995

This book is an introduction to computer graphics, now inevitably dated, but useful as an introduction to the topic. It covers visualization design, software selection, and a range of case studies. The appendices include a glossary of visualisation terms.

- Hierarchical and Geometrical Methods in Scientific Visualization, G Farin, B Hamann and H Hagen (eds), Springer, 2003
- Data Visualization Techniques, Chandrajit Bajaj (editor), John Wiley 1999

These books are largely of interest to programmers writing software that images volumes and shapes, and meshes, rather than end users of such software. The topics covered include: terrain modelling, multiresolution subdivision, wavelet-based scientific data compression, topology-based visualisation, data structures, data organisation and indexing schemes for scientific data visualisation.

- The Craft of Scientific Presentations, Michael Alley, Springer, New York, 2003

This text contains advice on making scientific presentations including the use of graphs and statistics in slides, especially how to improve PowerPoint presentations, and with short appendices that provide a checklist for scientific presentations and advice on the design of scientific posters.

C.2 Web-based resources

Some specific web sites with visualisation advice include the following :

Professor Antony Unwin: Professor Unwin holds the chair of computer-oriented statistics and data analysis at the University of Augsburg. His web pages are at <http://stats.math.uni-augsburg.de/~unwin/> and his main research topics include visualising information, large data sets, exploratory data analysis and exploratory modelling analysis, parallel coordinate systems, and software development. The web pages include much downloadable software, mainly written for the Macintosh operating system.

Visualization and Data Analysis Resource: these pages at <http://www.sdsc.edu/~johnson/projects/resource/catalog.html> provide an extensive collection of links to various visualisation-related material available on the web. They contain information on visualisation techniques, hardware and software systems, and academic and government laboratory research groups.

VRVis: VRVis Research Center in Vienna, Basic Research in Visualisation. The web pages of this group are at <http://www.vrvis.at/vis/>. There is a wide range of visualisation resources including links to reports of research work by students based at the Center. This group also does work on the parallel co-ordinate technique and the site includes parallel co-ordinates resources.

Professor Alfred Inselberg Professor Inselberg is the inventor of the parallel co-ordinate technique. His website, at <http://www.math.tau.ac.il/~%7Eaiisreal/>, includes the background to the technique, a list of his key publications in the area and a short tutorial on the technique (approached from the point of view of a geometrist).