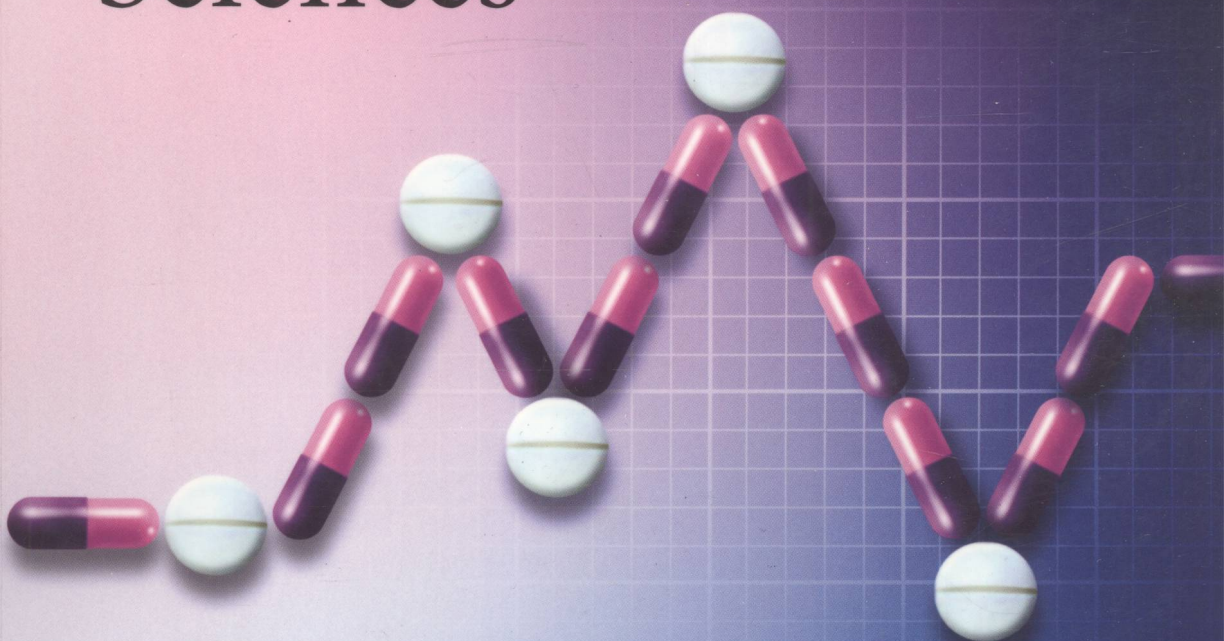


Philip Rowe

Essential Statistics

for the

Pharmaceutical Sciences



 WILEY



R911
R879

Essential Statistics for the Pharmaceutical Sciences

Philip Rowe

Liverpool John Moores University, UK



Copyright © 2007 John Wiley & Sons Ltd, The Atrium, Southern Gate, Chichester,
West Sussex PO19 8SQ, England

Telephone (+44) 1243 779777

Email (for orders and customer service enquiries): cs-books@wiley.co.uk
Visit our Home Page on www.wileyeurope.com or www.wiley.com

All Rights Reserved. No part of this publication may be reproduced, stored in a retrieval system or transmitted in any form or by any means, electronic, mechanical, photocopying, recording, scanning or otherwise, except under the terms of the Copyright, Designs and Patents Act 1988 or under the terms of a licence issued by the Copyright Licensing Agency Ltd, 90 Tottenham Court Road, London W1T 4LP, UK, without the permission in writing of the Publisher. Requests to the Publisher should be addressed to the Permissions Department, John Wiley & Sons Ltd, The Atrium, Southern Gate, Chichester, West Sussex PO19 8SQ, England, or emailed to permreq@wiley.co.uk, or faxed to (+44) 1243 770620.

Designations used by companies to distinguish their products are often claimed as trademarks. All brand names and product names used in this book are trade names, service marks, trademarks or registered trademarks of their respective owners. The Publisher is not associated with any product or vendor mentioned in this book.

This publication is designed to provide accurate and authoritative information in regard to the subject matter covered. It is sold on the understanding that the Publisher is not engaged in rendering professional services. If professional advice or other expert assistance is required, the services of a competent professional should be sought.

Other Wiley Editorial Offices

John Wiley & Sons Inc., 111 River Street, Hoboken, NJ 07030, USA
Jossey-Bass, 989 Market Street, San Francisco, CA 94103-1741, USA
Wiley-VCH Verlag GmbH, Boschstr. 12, D-69469 Weinheim, Germany
John Wiley & Sons Australia Ltd, 33 Park Road, Milton, Queensland 4064, Australia
John Wiley & Sons (Asia) Pte Ltd, 2 Clementi Loop #02-01, Jin Xing Distripark, Singapore 129809
John Wiley & Sons Canada Ltd, 6045 Freemont Blvd, Mississauga, Ontario, L5R 4J3

Wiley also publishes its books in a variety of electronic formats. Some content that appears in print may not be available in electronic books.

Anniversary Logo Design: Richard J. Pacifico

Library of Congress Cataloguing-in-Publication Data

Rowe, Philip.

Essential statistics for the pharmaceutical sciences / Philip Rowe.

p. ; cm.

Includes bibliographical references and index.

ISBN-13: 978-0-470-03470-5 (cloth : alk. paper)

ISBN-13: 978-0-470-03468-2 (pbk. : alk. paper)

1. Drugs—Research—Statistical methods. 2. Pharmacology—Statistical methods. I. Title.

[DNLM: 1. Pharmacology—methods. 2. Statistics. QV 20.5 R879e 2007]

RS57.R69 2007

615'.1072—dc22

2006102028

British Library Cataloguing in Publication Data

A catalogue record for this book is available from the British Library

ISBN-13: 978 0-470-03470-5 (HB)

ISBN-13: 978 0-470-03468-2 (PB)

Typeset in 10.5/12.5pt Minion by Thomson Digital

Printed in Great Britain by Antony Rowe Ltd., Chippenham, Wiltshire

This book is printed on acid-free paper responsibly manufactured from sustainable forestry in which at least two trees are planted for each one used for paper production.

Essential Statistics for the Pharmaceutical Sciences

To

Carol, Joshua and Nathan

Preface

At whom is this book aimed?

Statistics users rather than statisticians

As a subject, statistics is boring, irrelevant and incomprehensible. Well, traditionally it is anyway. The subject does not have to be that bad; It simply has not escaped from a self-imposed time-warp. Thirty years ago, there were no easy-to-use computer packages that would do all the hard sums for us. So, inevitably, the subject was heavily bogged down in the ‘How to’ questions – how do we most efficiently calculate the average weight of 500 potatoes? How do we work out whether the yield of turnips really did increase when we used a couple of extra shovels of horse manure? How should we calculate whether there really is a relationship between rainfall and the size of our apples? Thirty years ago, statistical authors had a genuine excuse for their books being clogged up with detailed methods of calculation. However, one might ask, what is their excuse today? Nowadays, nobody in their right mind would manually calculate a complex statistical test, so why do they still insist on telling us how to do it?

Clearly the world does need genuine ‘statisticians’ to maintain and improve analytical methods and such people do need to understand the detailed working of specific procedures. However, the majority of us are not statisticians as such; we just want to use statistical methods to achieve a particular end. We should be distinguishing between ‘statisticians’ and ‘statistics users’. It is rather like the distinction between automotive engineers and motorists. Automotive engineering is an honourable and necessary profession, but most of us are just simple motorists. Fortunately, within the field of motoring, the literature does respect the difference. If a book is written for motorists, it will simply tell us that, if we want to go faster, we should press the accelerator. It will not bore us stiff trying to explain how the fuel injection system achieves that particular end. Unfortunately that logic has never crept into statistics books. The wretched things still insist on trying to explain the internal workings of the Kolmogorov–Smirnov test to an audience who could not care less.

Well, good news! This book is quite happy to treat any statistical calculation as a black box. It will explain what needs to go into the box and it will explain what comes out the other end, but you can remain as ignorant about what goes on inside the box as you are about how your power steering works. This approach is not just lazy or negative. By stripping away all the irrelevant bits, we can focus on the aspects that

actually matter. This book will try to concentrate on those issues that statistics users really do need to understand:

- Why are statistical procedures necessary at all?
- How can statistics help in planning experiments?
- Which procedure should I employ to analyse the results?
- What do the statistical results actually mean when I have got them?

Who are these ‘statistics users’?

The people that this book is aimed at are those thousands of people who have to use statistical procedures without having any ambition to become statisticians. There are any number of student programmes, ranging from pharmacology and botany through to business studies and psychology, which will include an element of statistics. These students will have to learn to use at least the more basic statistical methods. There are also those of us engaged in research in academia or industry. Some of us will have to carry out our own statistical analyses and others will be able to call on the services of professional statisticians. However, even where professionals are to hand, there is still the problem of communication. If you do not even know what the words mean, you are going to have great difficulty explaining to a statistician exactly what you want to do. The intention is that all of the above should find this book useful.

If you are a statistics student or a professional statistician, my advice would be to put this book down now! You will probably find its dismissive attitude towards the mathematical basis of your trade extremely irritating. Indeed, if at least one traditional statistician does not complain bitterly about this book, I shall be rather disappointed.

To what subject area is the book relevant?

All the statistical procedures and tests mentioned are illustrated with practical examples and data sets. The cases are drawn from the pharmaceutical sciences and this is reflected in the book’s title. However, pretty well all the methods described and the principles explored are perfectly relevant to a wide range of scientific research, including pharmaceutical, biological, biomedical and chemical sciences.

At what level is it aimed?

The book is aimed at undergraduate science students and their teachers and less experienced researchers.

The early chapters (1–5) are fairly basic. They cover data description (mean, median, mode, standard deviation and quartile values) and introduce the problem of describing uncertainty due to sampling error (SEM and 95 per cent confidence interval for the mean). In theory, much of this should be familiar from secondary education, but in the author's experience, the reality is that many new students cannot (for example) calculate the median for a small data set. These chapters are therefore relevant to level 1 students, for either teaching or revision purposes.

Chapters 6–17 then cover the most commonly used statistical tests. Most undergraduate programmes will introduce such material fairly early on (levels 1 or 2). The approach used is not the traditional one of giving equal weight to a wide range of techniques. As the focus of the book is the various issues surrounding statistical testing rather than methods of calculation, one test (the two-sample *t*-test) has been used to illustrate all the relevant issues (Chapters 6–11). Further chapters (12–17) then deal with other tests more briefly, referring back to principles that have now been established.

The final chapters (18 and 19) cover some real-world problems that students probably will not run into until their final year, when carrying out their own independent research projects. The issues considered are multiple testing (all too common in student projects!) and questionnaire design and analysis. While the latter does not introduce many fundamentally new concepts, the use of questionnaires has increased so much, it seemed useful to bring together all the relevant points in a single resource.

Key point and pirate boxes

Key point boxes

Throughout the book you will find key point boxes that look like this:



Proportions of individuals within given ranges

For data that follow a normal distribution:

- about two-thirds of individuals will have values within 1 SD of the mean;
- about 95 per cent of individuals will have values within 2 SD of the mean.

These never provide new information. Their purpose is to summarize and emphasize key points.

Pirate boxes

You will also find pirate boxes that look like this:



Switch to a one-sided test after seeing the results

Even today, this is probably the best and most commonly used statistical fiddle.

You did the experiment and analysed the results by your usual two-sided test. The result fell just short of significance (P somewhere between 0.05 and 0.1) There is a simple solution, guaranteed to work every time. Re-run the analysis, but change to a one-sided test, testing for a change in whatever direction you now know the results actually suggest.

Until the main scientific journals get their act into gear, and start insisting that authors register their intentions in advance, there is no way to detect this excellent fiddle. You just need some plausible reason why you 'always intended' to do a one-tailed test in this particular direction, and you are guaranteed to get away with it.

These are written in the style of Machiavelli, but are not actually intended to encourage statistical abuse. The point is to make you alert for misuses that others may try to foist upon you. Forewarned is forearmed. The danger posed is reflected by the number of skull and cross-bone symbols.



Minor hazard. Abuse easy to spot or has limited potential to mislead.



Moderate hazard. The well-informed (e.g. readers of this book) should spot the attempted deception.



Severe hazard. An effective ruse that even the best informed may suspect, but never be able to prove.

Fictitious data

Throughout this book, all the experiments described are entirely fictitious as are their results. Generally the structure of the experiments and the results are realistic. In a few cases, both the structure of an experiment and consequently the analysis of the results

may be somewhat simpler than would often be seen in the real world. Clarity seems more important than strict realism.

At several points, judgements are quoted as to how greatly a measured end-point would need to change, for that change to be of any practical consequence. The values quoted are essentially arbitrary, being something that appears realistic to the author. Hopefully these are reasonable estimates, but none of them should be viewed as expert opinion.

A potted summary of this book

The book is aimed at those who have to use statistics, but have no ambition to become statisticians *per se*. It avoids getting bogged down in calculation methods and focuses instead on crucial issues that surround data generation and analysis (sample size estimation, interpretation of statistical results, the hazards of multiple testing, potential abuses, etc.). In this day of statistical packages, it is the latter that cause the real problems, not the number-crunching.

The book's illustrative examples are all taken from the pharmaceutical sciences, so students (and staff) in the areas of pharmacy, pharmacology and pharmaceutical science should feel at home with all the material. However, the issues considered are of concern in most scientific disciplines and should be perfectly clear to anybody from a similar discipline, even if the examples are not immediately familiar. Material is arranged in a developmental manner. Initial chapters are aimed at level 1 students; this material is fairly basic, with special emphasis on random sampling error. The next section then covers key concepts that may be introduced at levels 1 or 2. The final couple of chapters are most likely to be useful during final year research projects; These include one on questionnaire design and analysis.

The book is not tied to any specific statistical package. Instructions should allow readers to enter data into any package and find the key parts of the output. Specific instructions for performing all the procedures, using Minitab or SPSS, are provided in a linked web site (www.staff.ljmu.ac.uk/phaprowe/pharmstats.htm).

Statistical packages

There are any number of statistical packages available. It is not the intention of this book to recommend any particular package. Quite frankly, most of them are not worth recommending.

Microsoft XL

Probably the commonest way to collect data and perform simple manipulations is within a Microsoft XL spread-sheet. Consequently, the most obvious way to carry out statistical analyses of such data would seem to lie within XL itself. Let me give you my first piece of advice. Do not even consider it! The data analysis procedures within XL are rubbish – a very poor selection of procedures, badly implemented (apart from that, they are OK). If anybody in the Microsoft Corporation had a clue, they could have cleaned up in this area years ago. The opposition is pretty feeble and a decent statistical analysis package built into XL would have been unstoppable.

It is only at the most basic level that XL is of any real use (calculation of the mean, SD and SEM). It is therefore mentioned in some of the early chapters but not thereafter.

Other packages

A small survey by the author suggests that, in the subject areas mainly targetted by this book, only Minitab and SPSS are used with any frequency. Other packages seem to be restricted to very small numbers of departments. A decision was taken not to include blow-by-blow accounts of how to perform specific tests using any package, as this would excessively limit the book's audience. Instead, general comments are made about:

- entering data into packages;
- the information that will be required before any package can carry out the procedure;
- what to look for in the output that will be generated.

The last point is usually illustrated by generic output. This will not be in the same format as that from any specific package, but will present information that they should all provide.

Detailed instructions for Minitab and SPSS on the web site

As Minitab and SPSS clearly do have a significant user base, detailed instructions on how to use these packages to execute the procedures in this book, will be made available through the website associated with the book (www.staff.ljmu.ac.uk/phaprowe/pharmstats.htm). These cover how to:

- arrange the data for analysis;
- trigger the appropriate test;
- select appropriate options where relevant;
- find the essential parts of the output.

Contents

Preface	xiii
Statistical packages	xix
PART 1: DATA TYPES	1
1 Data types	3
1.1 Does it really matter?	3
1.2 Interval scale data	4
1.3 Ordinal scale data	4
1.4 Nominal scale data	5
1.5 Structure of this book	6
1.6 Chapter summary	6
PART 2: INTERVAL-SCALE DATA	7
2 Descriptive statistics	9
2.1 Summarizing data sets	9
2.2 Indicators of central tendency – mean, median and mode	10
2.3 Describing variability – standard deviation and coefficient of variation	16
2.4 Quartiles – another way to describe data	20
2.5 Using computer packages to generate descriptive statistics	23
2.6 Chapter summary	25
3 The normal distribution	27
3.1 What is a normal distribution?	27
3.2 Identifying data that are not normally distributed	28
3.3 Proportions of individuals within one or two standard deviations of the mean	31
3.4 Chapter summary	34
4 Sampling from populations – the SEM	35
4.1 Samples and populations	35
4.2 From sample to population	36
4.3 Types of sampling error	37
4.4 What factors control the extent of random sampling error?	39

4.5	Estimating likely sampling error – The SEM	42
4.6	Offsetting sample size against standard deviation	46
4.7	Chapter summary	46
5	Ninety-five per cent confidence interval for the mean	49
5.1	What is a confidence interval?	50
5.2	How wide should the interval be?	50
5.3	What do we mean by '95 per cent' confidence?	51
5.4	Calculating the interval width	52
5.5	A long series of samples and 95 per cent confidence intervals	53
5.6	How sensitive is the width of the confidence interval to changes in the SD, the sample size or the required level of confidence?	54
5.7	Two statements	56
5.8	One-sided 95 per cent confidence intervals	57
5.9	The 95 per cent confidence interval for the difference between two treatments	60
5.10	The need for data to follow a normal distribution and data transformation	61
5.11	Chapter summary	65
6	The two-sample t-test (1). Introducing hypothesis tests	67
6.1	The two sample t -test – an example of a hypothesis test	68
6.2	'Significance'	74
6.3	The risk of a false positive finding	75
6.4	What factors will influence whether or not we obtain a significant outcome?	76
6.5	Requirements for applying a two-sample t -test	79
6.6	Chapter summary	80
7	The two-sample t-test (2). The dreaded P value	83
7.1	Measuring how significant a result is	83
7.2	P values	84
7.3	Two ways to define significance?	85
7.4	Obtaining the P value	86
7.5	P values or 95 per cent confidence intervals?	86
7.6	Chapter summary	87
8	The two-sample t-test (3). False negatives, power and necessary sample sizes	89
8.1	What else could possibly go wrong?	90
8.2	Power	91
8.3	Calculating necessary sample size	94
8.4	Chapter summary	101
9	The two-sample t-test (4). Statistical significance, practical significance and equivalence	103
9.1	Practical significance – is the difference big enough to matter?	104

9.2	Equivalence testing	107
9.3	Non-inferiority testing	111
9.4	<i>P</i> values are less informative and can be positively misleading	113
9.5	Setting equivalence limits prior to experimentation	115
9.6	Chapter summary	116
10	The two-sample <i>t</i>-test (5). One-sided testing	117
10.1	Looking for a change in a specified direction	118
10.2	Protection against false positives	120
10.3	Temptation!	121
10.4	Using a computer package to carry out a one-sided test	125
10.5	Should one-sided tests be used more commonly?	126
10.5	Chapter summary	126
11	What does a statistically significant result really tell us?	127
11.1	Interpreting statistical significance	127
11.2	Starting from extreme scepticism	131
11.3	Chapter summary	132
12	The paired <i>t</i>-test – comparing two related sets of measurements	133
12.1	Paired data	133
12.2	We could analyse the data using a two-sample <i>t</i> -test	135
12.3	Using a paired <i>t</i> -test instead	135
12.4	Performing a paired <i>t</i> -test	136
12.5	What determines whether a paired <i>t</i> -test will be significant?	138
12.6	Greater power of the paired <i>t</i> -test	139
12.7	The paired <i>t</i> -test is only applicable to naturally paired data	139
12.8	Choice of experimental design	140
12.9	Requirements for applying a paired <i>t</i> -test	141
12.10	Sample sizes, practical significance and one-sided tests	141
12.11	Summarizing the differences between the paired and two-sample <i>t</i> -tests	143
12.12	Chapter summary	144
13	Analyses of variance – going beyond <i>t</i>-tests	145
13.1	Extending the complexity of experimental designs	146
13.2	One-way analysis of variance	146
13.3	Two-way analysis of variance	156
13.4	Multi-factorial experiments	164
13.5	Keep it simple – Keep it powerful	165
13.6	Chapter summary	167
14	Correlation and regression – relationships between measured values	169
14.1	Correlation analysis	170
14.2	Regression analysis	178

14.3	Multiple regression	185
14.4	Chapter summary	192
PART 3: NOMINAL-SCALE DATA		195
15	Describing categorized data	197
15.1	Descriptive statistics	198
15.2	Testing whether the population proportion might credibly be some pre-determined figure	202
15.3	Chapter summary	207
16	Comparing observed proportions – the contingency chi-square test	209
16.1	Using the contingency chi-square test to compare observed proportions	210
16.2	Obtaining a 95 per cent CI for the change in the proportion of expulsions – is the difference large enough to be of practical significance?	213
16.3	Larger tables – attendance at diabetic clinics	214
16.4	Planning experimental size	217
16.5	Chapter summary	219
PART 4: ORDINAL-SCALE DATA		221
17	Ordinal and non-normally distributed data. Transformations and non-parametric tests	223
17.1	Transforming data to a normal distribution	224
17.2	The Mann-Whitney test – a non-parametric method	228
17.3	Dealing with ordinal data	233
17.4	Other non-parametric methods	235
17.5	Chapter summary	242
	Appendix to chapter 17	242
PART 5: SOME CHALLENGES FROM THE REAL WORLD		245
18	Multiple testing	247
18.1	What is it and why is it a problem?	247
18.2	Where does multiple testing arise?	248
18.3	Methods to avoid false positives	250
18.4	The role of scientific journals	254
18.5	Chapter summary	255
19	Questionnaires	257
19.1	Is there anything special about questionnaires?	258
19.2	Types of questions	258
19.3	Designing a questionnaire	262
19.4	Sample sizes and return rates	263

19.5	Analysing the results	265
19.6	Confounded epidemiological data	266
19.7	Multiple testing with questionnaire data	271
19.8	Chapter summary	272
PART 6: CONCLUSIONS		275
20	Conclusions	277
20.1	Be clear about the purpose of the experiment	277
20.2	Keep the experimental design simple and therefore clear and powerful	278
20.3	Draw up a statistical analysis plan as part of the experimental design – it is not a last minute add-on	279
20.4	Explore your data visually before launching into statistical testing	280
20.5	Beware of multiple analyses	281
20.6	Interpret both significance and non-significance with care	282
Index		283