

Boekbesprekingen

Redactie

Alex J. Koning

postadres: Econometrisch Instituut
Erasmus Universiteit Rotterdam
Postbus 1738
3000 DR Rotterdam
telefoon: 010-4081268/2231
fax: 010-4527746
e-mail: koning@few.eur.nl

Besproken boeken in Kwantitatieve Methoden nr. 59

Coelli, T., Prasada Rao, D.S., Battese, G.E.
An introduction to efficiency and productivity analysis

Christensen, R.
Log-linear models and logistic regression

Mandelbrot, B.B.
Fractals and scaling in finance. Discontinuity, concentration, risk

Krause, A., Olson, M.
The basics of S and S-Plus

Devlin, B., Fienberg, S.E, Resnick, D.P., Roeder, K.
Intelligence, genes and success. Scientists respond to the Bell curve

Mann, P.S.
Introductory statistics

Zeisel, H., Kaye, D.
Prove it with figures. Empirical methods in law and litigation

Kotz, S., Johnson, N.L.
Breakthroughs in statistics

Buldygin, V., Solntsev, S.
Asymptotic behaviour of linearly transformed sums of random variables

Lange, K.
Mathematical and statistical methods for genetic analysis

Lindsey, J.K.

Applying generalized linear models

Klein, J.P., Moeschberger, M.L.

Survival analysis. Techniques for censored and truncated data

Gerber, S., Voelkl, K.E.

The SPSS guide to The new statistical analysis of data

Wilcox, R.R.

Introduction to robust estimation and hypothesis testing

Coelli, T., Prasada Rao, D.S., Battese, G.E.

An introduction to efficiency and productivity analysis

1997, Kluwer, Dordrecht,

296 blz, fl 215.00, ISBN 0-7923-8060-6.

Dit boek is bedoeld als "first port of call" voor ieder die zich bezig houdt of wil gaan houden met de meting van de "performance" van bedrijven (in de breedste zin van het woord genomen), zowel cross-sectioneel als intertemporeel. Anders gezegd, het gaat om het meten en vergelijken van efficiency en productiviteitsveranderingen. Een bedrijf heet inefficiënt als het bij gelijkblijvende hoeveelheden input zijn productie kan vergroten, of als het bij gelijkblijvende hoeveelheden output kan bezuinigen op zijn verbruik. Om de mate van (in-)efficiency te kunnen meten is kennis van de productiestructuur ofwel de technologie nodig. In het algemeen kan deze kennis alleen verkregen worden door econometrische technieken los te laten op geschikte data, bij voorkeur bedrijfsmatige micro-data.

Er is sprake van productiviteitsverandering als het verloop in de tijd van de geproduceerde hoeveelheden verschilt van dat van de verbruikte hoeveelheden. Een belangrijke vraag is die naar de mogelijkheid om gemeten productiviteitsverandering te ontbinden in betekenisvolle componenten zoals verandering in efficiency en technische verandering. Ook dit kan niet zonder econometrie en data.

Naar ik meen, zijn de auteurs in hun opzet bijzonder goed geslaagd. Het boek bevat een goede mix van theorie en toepassing, en kan op diverse kennisniveaus met profijt gelezen en gebruikt worden. De relevante theorie wordt zorgvuldig besproken, met verwijzingen naar de literatuur voor ieder die op verdieping uit is. De kracht van het boek ligt in de uitgewerkte voorbeelden en toepassingen, die zo duidelijk zijn dat ieder, gewapend met de in het boek besproken softwarepakketten, onmiddellijk aan de slag kan.

In de hoofdstukken 2 en 3 wordt de traditionele aanpak besproken: het schatten van productiefuncties met behulp van de kleinste kwadraten methode en het op basis daarvan meten van technische verandering. Ook komt het schatten van kosten-, winst-, en afstandsfuncties aan de orde.

In de hoofdstukken 4 en 5 gaat het over het berekenen van totale factorproductiviteits-indexcijfers. Daartoe wordt niet alleen de axiomatische indextheorie besproken, maar ook de economische. Het wordt daardoor duidelijk welke veronderstellingen aan deze methode ten grondslag liggen.

De hoofdstukken 6 en 7 gaan over het meten van efficiency met behulp van niet-parametrische data-omhullingsmethoden (Data Envelopment Analysis).

In de hoofdstukken 8 en 9 wordt de stochastische, parametrische "frontier" methode voor het berekenen van efficiency-maatstaven besproken.

Hoofdstuk 10 legt dan vervolgens uit hoe men met de DEA- en de "frontier"-methode maatstaven voor technische verandering en verandering in efficiency kan berekenen.

Globaal gezien, worden dus aan elk van de vier methoden twee hoofdstukken gewijd. Het eerste hoofdstuk van elk paar bespreekt de grondslagen van theorie en toepassing, terwijl het tweede uitwerkingen, verfijningen en alternatieven aanbiedt. Dit maakt het boek zeer geschikt voor praktisch ingestelde lezers. Men behoeft niet alles te lezen om met een bepaalde methode aan het werk te kunnen gaan. Het laatste, elfde, hoofdstuk zet bovendien de sterke en zwakke punten van de vier methoden nog eens overzichtelijk bij elkaar en geeft nuttige aanwijzingen voor de toepassing.

De Appendix geeft nadere uitleg bij de in het boek prominent figurerende, door Coelli ontworpen, softwarepakketten, te weten FRONTIER 4.1 (voor stochastische frontiersanalyse), DEAP 2.1 (voor data-omhulling), en TFP 1.0 (voor productiviteits-indexcijfers). Een kind, zoals deze recensent, kan hiermee de was doen.

Uiteraard zou men zo hier en daar met de auteurs in discussie willen treden. Zelf zou ik die noodzaak voelen bij de bespreking van het onderwerp schaal (returns to scale) en schaal-efficiency, alsmede bij de aan de methode der indexcijfers toegeschreven veronderstellingen. Maar dit zijn ondergeschikte punten.

Aan de uitvoering van het boek is door de auteurs veel zorg besteed. Op een enkel punt ziet men dat het een co-productie is. De in hoofdstuk 5 opgevoerde Malmquist productiviteits-index wordt in hoofdstuk 10 ineens "Malmquist TFP (totale factorproductiviteits-) index" genoemd. Typefouten vond ik weinig. Opmerkelijk was slechts dat het door deze recensent ontvangen exemplaar verkeerd-om ingebonden was.

Veel lof dus voor dit boek. Mijns inziens verdient het een plaats in het universitaire economie curriculum.

dr B.M. Balk
Department of Statistical Methods
Statistics Netherlands

Christensen, R.

Log-linear models and logistic regression

1997, Springer texts in statistics,
 Springer-Verlag, Berlin,
 495 blz, DM 108.00, ISBN 0-387-98247-7.

Dit boek, in de serie "Springer texts in statistics", is een tweede editie; de eerste editie, getiteld "Log-linear models", dateert van 1990. De belangrijkste verandering is dat in deze nieuwe editie aanmerkelijk meer aandacht wordt besteed aan logistische regressie. Voorts wordt aandacht besteed aan de theorie van gegeneraliseerde lineaire modellen - waarvan het log-lineaire en het logistische model uiteraard speciale gevallen zijn. Ook het gebruik van grafentheorie en matrix-algebra bij de analyse van log-lineaire modellen wordt behandeld. Tenslotte is er een hoofdstuk toegevoegd over Bayesiaanse binomiale regressie. Het zwaartepunt van het boek ligt op het log-lineaire model voor de analyse van (meer-weg) kruistabellen.

De analogie tussen lineaire modellen (variantie-analyse en lineaire regressie) enerzijds en log-lineaire en logistische modellen anderzijds loopt als een rode draad door het boek; hierbij worden uiteraard ook de soms wezenlijke verschillen tussen deze modellen gesignaleerd. Een gedegen kennis van variantie-analyse en lineaire regressie ("at the masters degree level") wordt dan ook bij de lezer voorondersteld. In de latere hoofdstukken van het boek is tevens een degelijke basiskennis van de mathematische statistiek vereist, inclusief onderwerpen als likelihood theorie en matrix-algebra. Het boek lijkt in de eerste plaats te zijn geschreven voor statistici en gevorderde studenten in de (mathematische of toegepaste) statistiek, maar het kan ook zeker interessant zijn voor onderzoekers in vakgebieden waarin de statistiek een centrale rol speelt.

De eerste hoofdstukken van dit boek vertonen een duidelijke gradiënt in moeilijkheidsgraad. Het begint op zeer elementair niveau: zelfs de introductie van een simpel concept als de odds ratio neemt meerdere bladzijden in beslag, compleet met uitgewerkte voorbeelden. Naarmate ingewikkelder onderwerpen aan bod komen wordt echter al snel een grotere intellectuele inspanning van de lezer gevraagd. De relatief onervaren lezer kan uit de eerste helft van het boek een goede basiskennis verwerven; de resterende hoofdstukken fungeren dan als "capita selecta" waaruit eventueel later nog naar behoefte kan worden geput. Voor de gevorderde lezer zullen juist deze latere hoofdstukken het interessantste gedeelte van het boek vormen; de meer elementaire hoofdstukken kunnen hierbij dienen als overzichtelijke recapitulatie van grotendeels bekende theorie.

Als epidemioloog-statisticus constateerde uw recensent met interesse dat ook epidemiologische toepassingen beknopt aan bod komen. Helaas laat de auteur bij de behandeling van dit interessante toepassingsgebied nogal wat steken vallen. Om te beginnen zijn de definities van prospectief onderzoek en cohort-onderzoek al wat eigenaardig: er is hier voortdurend sprake van het trekken van aselechte steekproeven en dan kijken wie er (bijv.) een hartinfarct hebben doorgemaakt ... in het verleden (?). Vervolgens lezen wij de obligate waarschuwing tegen het fenomeen "confounding" (in dit boek aangeduid als "Simpson's paradox"): in niet-experimenteel onderzoek hoeft een gevonden associatie niet noodzakelijkerwijs op een causaal verband te berusten, het kan immers zijn dat de echte oorzaak een andere factor is die met de bestudeerde variabelen samenhangt. Na deze gerechtvaardigde waarschuwing wordt de lezer echter finaal op het verkeerde been gezet met de mededeling: "We hope that random sampling within the population of people minimizes these effects" (p 119). Dat is dan ijdele hoop! Niet aselechte steekproeftrekking maar multiple regressie-analyse (met de veronderstelde "confounders" als covariabelen in het model) is een universeel toepasbare methode voor eliminatie van "confounding". Een akelige uitglijder, uitgerekend in een boek over regressiemodellen!

Niettenstaande deze kritiek moet worden vastgesteld dat dit boek toch veel te bieden heeft. De theorie wordt systematisch opgebouwd, en de auteur heeft zich duidelijk ingespannen om een student-vriendelijke didactische opbouw te realiseren. De theorie wordt verduidelijkt met vele uitgewerkte voorbeelden. Hierbij worden datasets gebruikt uit toepassingsgebieden, uiteenlopend van ruimtevaarttechnologie en genees-

middelenonderzoek tot de cognitieve vermogens van studenten en zeeleeuwen. Hierbij wordt ook het gebruik van moderne software gedemonstreerd (voornamelijk SAS en BMDP, ook wel GLIM). Aan het eind van ieder hoofdstuk treffen we nog enkele oefeningen, waaraan de lezer kan toetsen of hij/zij de stof beheerst (uitwerkingen van deze oefeningen zijn echter niet in het boek opgenomen). De appendix bevat, naast een tabel van de chi-kwadraat verdeling, het complete Griekse alfabet.

Al met al kan dit boek worden aanbevolen aan de bovengenoemde doelgroepen. Rest nog te vermelden dat het boek mooi is uitgegeven: het is een gebonden editie, op goed papier gedrukt in een prettig leesbaar lettertype. De prijs is naar hedendaagse maatstaven schappelijk.

P.J. Kostense

*Vakgroep Epidemiologie & Biostatistiek
Vrije Universiteit Amsterdam*

Mandelbrot, B.B.

Fractals and scaling in finance. Discontinuity, concentration, risk

1997, Springer-Verlag, Berlin,
455 blz, DM 84.00, ISBN 0-387-98363-5.

Readership: students and researchers interested in finance and modelling of financial (and other) data.

The book consists of five parts. In Part I (Chapters E1 – E4) the author gives a non-mathematical introduction to several terms such as 'scaling', 'Brownian Motion', 'wild randomness', 'Fractal Brownian Motion', etc. which appear throughout the book. The author explains that the 'models' he obtains are NOT to be confused with sophisticated economical or mathematical expressions. In his words "an ideal model is one that produces data streams that are hard to distinguish from actual records".

In Parts II – V several of the favourite topics of Mandelbrot are discussed in some length. Among others in Chapter E5 the author distinguishes several types of randomness. To this end he starts from i.i.d. random variables $U, U(1), U(2), \dots, U(n)$ and uses partial sums $S(n) = U(1) + U(2) + \dots + U(n)$ and partial maxima $M(n) = \max(U(1), \dots, U(n))$. The r.v. U is even portioning (EP) if all $U(i)$ have a similar contribution to the sum $S(n)$. U is concentrated portioning (CP) if $M(n)$ dominates the sum $S(n)$. Depending on whether n is small (Short Term ST) or large (Long term LT) we have mild randomness if U shows EP in ST and in LT. We have slow randomness if U shows EP in LT and CP in ST. Finally we have wild randomness if U shows CP in ST and CP in LT. Apart from this classification also other classifications and refinements are discussed in the book.

It is a pity that the author seems to be unaware of the huge literature on the so-called subexponential distributions functions and subexponential densities. This would have been useful in, for example, the discussion in Chapter E5 of the present book.

A second favourite topic is the notion of scaling. In Chapter E8 the author lists several interesting possible applications of scaling distributions. These are distributions of the form $\log(P\{U > u\}) = c - \beta \log(u)$. Scaling appears in a natural way in modelling financial data as follows. Let $L(t, T) = \log(Z(t+T)) - \log(Z(t))$ denote the return of an asset, where $Z(x)$ denotes the price of the asset at time x . There seems to be some empirical evidence that assets have a return with the characteristics of a scaling random variable i.e. the logarithm of the tail distribution of $L(t, T)$ is of the form $c(T) - \beta \log(u)$.

Going from scaling distributions to distributions with a regularly varying tail or in a domain of attraction is only as small step. The author shows that one can use L-stable (for example in Chapters E14-E16) distributions to model financial data. Empirical evidence shows that the L-stable market hypothesis is more consistent with the data than the Gaussian market hypothesis. Also the Fractional Brownian Motion (FBM) seems to yield good results. Recall that FBM is a random process $B(t)$ with Gaussian increments $DB(t, T) := B(t+T) - B(t)$ that satisfy $E(DB(t, T)) = 0$ and $\text{Var}(DB(t, T)) = T^{2H}$.

If the Hurst coefficient H equals $H = 1/2$, the process is the Wiener Brownian Motion. If $0 < H < 1/2$ or $1/2 < H < 1$ the process is neither a martingale nor a Markov process. FBM is discussed in Chapters E6. In chapters E19 and E20 the author concludes the book with some comments on the use and the limitations of martingale models.

In this brief summary of the book, I have limited myself only to some of the topics of the book. Part I gives indeed an interesting introduction and overview of randomness and scaling and can be easily read by an interested reader. Unfortunately most chapters of Parts II–V of the book are reproductions of articles which have been published (and refereed) in the period 1960-1996. The articles or revised versions of the articles are reprinted together with some additional comments and some critical remarks. As a summary, one of the merits of the book is that it gives a good overview of the works of Mandelbrot. Mandelbrot has original ideas and the recognition of the fractal nature of e.g. price variations has indeed influenced many researchers. In the book there are some new ideas and materials but there are so many overlaps with the existing papers and so many cross-references that it is only possible to read the book by continuously moving from one chapter to another chapter. If the reader wants to complete his or her library of Mandelbrot, (s)he should not hesitate to do so. If the reader is rather interested in results, it is perhaps better to consult and read some (or all) of the earlier published papers.

prof. dr E. Omey
EHSAL

Krause, A., Olson, M.

The basics of S and S-Plus

1997, Statistics and computing,
Springer-Verlag, Berlin,
270 blz, DM 48.00, ISBN 0-387-94985-2.

De groeiende populariteit en commercialisering van S-PLUS gaat gepaard met een groeiende stroom boeken over S-PLUS. De meerderheid hiervan valt te karakteriseren als uitgebreid, diepgravend en geschreven voor de 'computer literate'. Het boek van Krause en Olson verschilt wezenlijk met een oriëntatie op de lezer die een begin wil maken met het gebruik van de taal van S-PLUS, zowel interactief als voor het schrijven van eigen functies, en die daarbij op weg geholpen wil worden.

Het boek begint met vier inleidende hoofdstukken. Hierin wordt de lezer vertrouwd gemaakt met de geschiedenis van S-PLUS, de (menu) interface van S-PLUS (althans voor versie 4), en het gebruik van S-PLUS middels de commando taal. Dat laatste vindt plaats middels een tweetal interactieve sessies waarbij en passant een aantal nuttige functies en de belangrijkste data structuren worden geïntroduceerd.

De kern van het boek wordt gevormd door vier hoofdstukken over achtereenvolgens: graphics, beschrijvende statistiek, statistische modellen en programmeren. Het hoofdstuk over graphics behandelt in essentie alleen het weergeven van bivariate gegevens. Daarbij komt echter een veelheid aan onderwerpen aan de orde zoals plotsymbolen, titels, assen, grootte van de elementen in het plaatje en meerdere plots binnen een plaatje.

Verdere mogelijkheden voor de graphische weergave van zowel univariate als multivariate data worden geïllustreerd in het hoofdstuk over beschrijvende statistiek: met o.a. histogrammen, boxplots, maar ook trellis displays. Tevens vindt men hier enkele commando's voor het karakteriseren van data middels kengetallen. Het hoofdstuk over beschrijvende statistiek wordt besloten met enkele vluchtige hoofdstukken over (theoretische) kansverdelingen, het toetsen van hypothesen en ontbrekende waarden.

De statistische modellen van S-PLUS worden toegelicht aan de hand van lineaire regressie, ANOVA, gegeneraliseerde lineaire modellen en levensduur analyse. Er wordt een toelichting gegeven op de standaard S-PLUS manier voor het specificeren van de modellen via een model formule. Voorts vindt men een (gering en tamelijk *ad hoc* gekozen) aantal functies die van belang zijn bij het analyseren van data middels de genoemde statistische modellen.

In het hoofdstuk over programmeren bestaat uit een wat onevenwichtige mix van onderwerpen: van te verwachten onderwerpen als iteratie, functies, en debug hulpmiddelen, naar zeer eenvoudige zaken als het gebruik van de functie 'cat', tot zeer geavanceerde zaken als het methode/klasse mechanisme van S-PLUS voor object georiënteerd programmeren.

Het boek eindigt met een aantal restjes: het lezen en schrijven van data, hints en speciale onderwerpen. Het is duidelijk dat de auteurs door schade en schande wijs geworden zijn met betrekking tot het gebruik van S-PLUS. In deze hoofdstukken geven

zij een aantal vuistregels voor het efficiënt gebruik van S-PLUS die zij hebben ontwikkeld tijdens hun beginnersstrijd met S-PLUS. Behandeld worden onder andere het ontwikkelen van functies, het gebruik van verschillende directoriës, S-PLUS plaatjes en tekstverwerking (latex en word), bibliotheken van S-PLUS functies en het aanroepen van C en Fortran functies binnen S-PLUS. Erg nuttig ook is de uitgebreide verwijzing naar elektronische bronnen van informatie, zoals bijvoorbeeld het zeer actieve S-NEWS en de collectie van S-PLUS functies op de statlib server.

Het is een boek dat erop gericht is om een beginnende gebruiker bekend te maken met de commando-taal van S-PLUS, zowel voor het interactieve gebruik als voor het schrijven van functies. In veel eenvoudige voorbeelden worden belangrijke functies helder uitgelegd en de opgaven met uitgewerkte antwoorden stimuleren creatief gebruik. De eenzijdige nadruk op de computer kant met eenvoudige voorbeelden heeft echter ook zijn nadelen: statistisch en data-analytisch valt er weinig lol te beleven aan het boek. Het is daarmee geen boek dat de statisticus motiveert tot het gebruik van S-PLUS. De reeds gemotiveerde statisticus daarentegen krijgt een gedegen, behoudzame en goed leesbare inleiding tot de geavanceerde mogelijkheden die S-PLUS biedt middels de commando-taal.

dr D. Denteneer

Philips Research Laboratories Eindhoven

Devlin, B., Fienberg, S.E, Resnick, D.P., Roeder, K.

Intelligence, genes and success. Scientists respond to the Bell curve

1997, Springer-Verlag, Berlin,

385 blz, DM 49.00, ISBN 0-387-94986-0.

Het boek dat in de de sociale wetenschappen het meeste stof heeft doen opwaaien gedurende de afgelopen jaren is zonder twijfel *The Bell Curve* door Herrnstein en Murray. De auteurs van dit boek zijn in de media afgeschilderd als 'respectable wetenschappers die op verantwoorde wijze een belangrijk probleem bespreken' en als 'semi-wetenschappelijke racisten'. De meeste controversie betrof een bespreking van het verschil tussen de IQ-verdeling van Kaukasische Amerikanen en die van Afro-Amerikanen. De hypothese dat IQ in hoge mate erfelijk is wordt met kracht verdedigd, alsmede de consequenties die dit heeft voor het beleid (waarom zou de overheid in onderwijs investeren als mensen er niet veel slimmer van worden?). Volgens vele wetenschappers was er van alles mis met de empirische analyses in het boek; tijdens een lezingenmiddag georganiseerd door de Economische Sectie van de VVS 1995 bleek dat bijvoorbeeld de resultaten betreffende de relatie tussen opleiding en inkomen niet wordt gevonden met andere data of een andere specificatie (zie Hartog en Oosterbeek (1997)).

Deze bundel met reacties op *The Bell Curve* bestaat uit vijf delen:

1. Overview
2. The Genetics-Intelligence Link

3. Intelligence and the Measurement of IQ
4. Intelligence and Success: Reanalyses of Data of the NLSY
5. *The Bell Curve* and Public Policy

Dit boek is niet de eerste reactie op *The Bell Curve*. Er zijn talloze boekbesprekingen verschenen (niet alleen in kranten, maar ook in tijdschriften als de *Scientific American* en de *Journal of the American Statistical Association*). Een groot aantal boekbesprekingen beperkten zich niet alleen tot het bespreken van het boek, maar gingen uitgebreid in op alternatieve interpretaties van het geboden materiaal. In het voorwoord van Devlin *et al.* worden vijf (*sic*) boeken genoemd die reageren op *The Bell Curve*.

Wat is de toegevoegde waarde van de bundel van Devlin *et al.*? In tegenstelling tot (sommige) andere commentaren wordt hier diep en vanuit een wetenschappelijk perspectief op de zaak ingegaan. Inderdaad: er worden door verschillende auteurs van naam en faam nieuwe empirische en theoretische analyses gepresenteerd die aantonen dat de gegevens die Herrnstein en Murray gebruiken meerdere interpretaties toelaten, en dat de gebruikte statistische methodologie soms zeker voor verbetering vatbaar is. In het laatste hoofdstuk vatten Resnick en Fienberg de bevindingen samen en, misschien niet geheel verbazingwekkend, Herrnstein en Murray komen er slecht van af. Er zijn allerlei kritieken aan te voeren op de analyses. Deze kritieken komen uit allerlei wetenschapsgebieden: de statistiek, genetica, sociologie, psychologie, etc.

Voor wie is dit boek nu bedoeld? Ik zou het eigenlijk niet weten. In de eerste twee hoofdstukken worden de bevindingen uit *The Bell Curve* samengevat en iemand die het origineel te omvangrijk vindt zou die kunnen lezen. De bijdragen worden door de redacteurs zelf aangeduid als papers die niet zelfstandig in de vakliteratuur een plaats zullen vinden. Ik denk dat het boek voor mensen, die geïnteresseerd zijn in het problemen die worden besproken in *The Bell Curve*, interessant kan zijn als achtergrondliteratuur. Een omvangrijke reactie op het werk van iemand anders is voor een argeloze lezer echter zelden boeiend.

Verwijzingen:

- [1] J. Hartog en H. Oosterbeek (1997), Health, Wealth, and Happiness, Tinbergen Institute, Discussion Paper TI 97-034/3.
- [2] R.J. Herrnstein en C. Murray (1994) *The Bell Curve: Intelligence and Class Structure in American Life* New York: The Free Press.

dr R.H. Koning
 Vakgroep Econometrie
 Vrije Universiteit Amsterdam

Mann, P.S.

Introductory statistics

1998, John Wiley & Sons Ltd., Chichester,
xxiii+789 blz, £ 22.95, ISBN 0-471-16546-8.

Introductory Statistics is exactly what the title suggest. An introduction into statistics for students who have no strong background in mathematics. It is well written and uses many educational devices to help the beginning student. Each chapter is broken down into short, comprehensive subsections. Each concept is introduced carefully, and followed by many examples and exercises. Each chapter ends with a comprehensive glossary of key terms, a list of key formulas, and a self-review test. Also each chapter contains computer exercises and examples using MINITAB. The book includes a comprehensive glossary of symbols and an overview of key formulas ordered by chapter (topic). It also provides the student with the answers to all odd numbered exercises and self-review tests.

I enjoyed the realistic case studies and the 'observations' and 'warnings' in the text. The case studies come from different fields in the social and medical sciences and are real life examples (sources are quoted). Although real life, the author has made a careful selection and all examples are transparent and easy to grasp. The author clearly has many years of teaching and experience and knows where students may experience difficulties. Under the subheading 'observations' short but clear explanations are given on points students often wonder about. The 'warnings' warn students to avoid common mistakes. As an example, in the section on measures of dispersion for ungrouped data, subsection variance and standard deviation one finds the observation that the values of variance and standard deviation are never negative, and a clear explanation. One also finds the warning that $\sum x^2$ is not the same as $(\sum x)^2$, and again an explanation why.

In thirteen chapters the basics of statistics are explained. Chapter 1 gives an introduction in statistical concepts (e.g., descriptive vs inferential statistics, quantitative and qualitative variables, sample and population). Chapter 2 is about organizing data and describes frequency distributions and graphical representations (bar graphs, pie charts, stem and leaf plots). Chapter 3 discusses measures of central tendency and dispersion, and includes some excellent examples making sure that the students grasps the concepts. Chapter 4 introduces probability and illustrates it with some good examples and fun case studies. Chapter 5 and 6 explain discrete and continuous random variables and the binomial, Poisson and normal distribution. Chapter 7 extends the concept of probability to that of a sampling statistic and introduces sampling distributions. Together with Appendix A on sampling surveys and design of experiments, it gives an excellent introduction into what readers of articles and journalists should know about surveys. Chapters 8 and 9 introduce the student to inferential statistics (estimation, confidence intervals, and hypothesis testing). The next three chapters introduce basic statistical tests. Chapter 9 is about hypothesis testing procedures for differences between two populations (means and proportions), and includes lifelike examples (e.g. estimation of poverty rates, clinical trials to test the effectiveness of

nicotine patch therapy for stopping smoking, differences in salaries between men and women). Chapter 11 thoroughly treats chi-square tests (e.g., goodness-of-fit test, homogeneity test), and chapter 12 introduces one way analysis of variance. The last chapter (13) explains correlation coefficients and introduces regression. This is where the book ends, multivariate analysis is not included.

The emphasis in Introductory Statistics is on applied statistics, which at the same time is the strength and weakness of the book. It focusses on comprehension and applicability of statistics. It does not discuss the mathematical fundament of statistics. It is not a book for 'math-stats', but it is an excellent introduction for novices in statistics, taking away most of the fear for statistics and introducing some fun reading the case studies. The many examples and exercises use realistic situations and after completing this book a student will be able to read and comprehend research reports and articles in applied journals in the social and medical sciences.

For introductory purposes, the book is a good buy for 23 dollars. Extra supplements (e.g., instructors manual, student study guide, and a test bank with multiple choice items in print and on diskette) are available. The data for the computer exercises are available on diskette and on the Wiley website. A disadvantage is that the computer exercises and instructions are only for MINITAB, not for SPSS or SAS. An inventive instructor can of course 'translate' these to SPSS, especially as the data are available both in MINITAB and ASCII-format. A second disadvantage for Dutch students is that the language is English.

This is a review of the third edition; the first edition was reviewed by Van Peet in KM (1996). Compared with the first edition, the examples, case studies and exercises are updated. Extra examples have been added to key concepts to enhance clarity and the MINITAB sections are updated to MINITAB for WINDOWS. Also exercises are added for graphing calculators (e.g. Texas Instruments TI-83), the programs needed for running complex statistical tests can be again downloaded from the Wiley website (<http://www.wiley.com.college/mann>). Also available on the website are suggestions for several one semester courses based on the book. In short it is a good educational tool. The main criticisms of Van Peet, however, remain unaddressed. Non-parametric statistics (except chi-square statistics) is not included. Analysis of variance is limited to one-way analysis of variance and regression analysis is limited to two variables only.

*dr E.D. de Leeuw
Methodika Amsterdam*

Zeisel, H., Kaye, D.

Prove it with figures. Empirical methods in law and litigation

1997, Statistics for social science and public policy,

Springer-Verlag, Berlin,

400 blz, DM 112.00, ISBN 0-387-94892-9.

Prove it with Figures is een boek dat met name is geschreven voor juristen. Het beoogt studenten, docenten, en advocaten een elementair begrip van het doel, de structuur, en het begrippenkader van sociaal-wetenschappelijk onderzoek bij te brengen, "so that, when the time comes, the lawyer can discuss it intelligently with an expert" (p. x). Het boek, dat voor een belangrijk deel is gebaseerd op een manuscript nagelaten door de inmiddels overleden Zeisel, behandelt achtereenvolgens het belang van empirisch onderzoek voor het bepalen van (juridische) causaliteit (hoofdstuk 1); de waarde en beperkingen van gecontroleerde experimenten (hoofdstuk 2), observatie-onderzoek (hoofdstuk 3), en epidemiologische studies (hoofdstuk 4). In hoofdstuk 5 worden twee technieken beschreven - triangulatie en replicatie - die kunnen worden gebruikt om de beperkingen van observatie-onderzoek (deels) te ondervangen. Hoofdstuk 6 handelt over P-waarden, significantie en standaarddeviaties, waarna in hoofdstuk 7 en 8 - voornamelijk door middel van voorbeelden - wordt ingaan op respectievelijk de basis-methodologie van steekproeven en inhoudsanalyse. In de daarop volgende hoofdstukken wordt het nut van sociaalwetenschappelijk onderzoek voor juridische doeleinden nader belicht aan de hand van een drietal onderwerpen; verzoeken tot gedingverplaatsing ("motions for change of venue"; hoofdstuk 9), merkenrecht (hoofdstuk 10 en 11), en juryselectie (hoofdstuk 12). In hoofdstuk 13 tenslotte volgt een verhandeling van de juridische en statistische vraagstukken met betrekking tot DNA-onderzoek en daarop gebaseerde (statistische) conclusies. Alle hoofdstukken zijn voorzien van een korte literatuurlijst. Sommige hoofdstukken zijn daarnaast voorzien van een lijst van kritische vragen met betrekking tot sociaal-wetenschappelijk onderzoek en/of een samenvatting. Ook bevat het boek een overzicht van termen die veelvuldig worden gebruikt in (de beschrijving van) sociaal-wetenschappelijk onderzoek.

Of het boek de kloof tussen juristen en statistici daadwerkelijk overbruggt, zoals de bedoeling was, is echter zeer de vraag. De eerste vijf hoofdstukken zijn zonder meer interessant en toegankelijk geschreven, terwijl de actuele voorbeelden die daarbij worden gegeven een goed inzicht geven in het belang en de beperkingen van sociaal-wetenschappelijk onderzoek voor juridische doeleinden. Zo behandelt hoofdstuk 4 onder andere sociaal-wetenschappelijk onderzoek naar de ziektes die (mogelijk) gerelateerd zijn aan blootstelling aan "Agent Orange" (een ontbladeringsmiddel dat door het Amerikaanse leger werd gebruikt in Vietnam), asbest en elektromagnetische velden. Tevens komen hier (de beperkingen van) studies aan de orde die betrekking hebben op een mogelijk causaal verband tussen borst implantaten en "connective tissue disorders" en het gebruik van Bendectin - een medicijn tegen misselijkheid - door zwangere Amerikaanse vrouwen en kinderen geboren met gereduceerde ledematen. Elk van deze onderwerpen heeft in Amerika tot omvangrijke rechtszaken geleid. Ook

in Nederland zijn deze onderwerpen niet geheel onbekend (vergelijk de zaak tegen de producenten van het medicijn Des), zij het dat de meeste onderwerpen hier (nog) niet tot uitgebreide gedingvoering aanleiding hebben gegeven.

Hoofdstuk 6 is van een geheel andere orde. In dit hoofdstuk, dat in feite de kern van het boek vormt, wordt een beroep gedaan op het vermogen van de jurist om P-waarden, statistisch significante waarden, betrouwbaarheidsintervallen en andere statistische berekeningen te doorgronden. Dat is voor de jurist geen eenvoudige taak, maar wel een uitermate relevante voor een goed begrip van sociaal- wetenschappelijk onderzoek. Merkwaardig is echter dat in een groot deel van de daarop volgende hoofdstukken betrekkelijk zelden op die kennis een beroep wordt gedaan. Dit wekt bij de jurist die verwoede pogingen heeft gedaan een en ander te begrijpen de nodige verbazing (en deels ook irritatie). In plaats daarvan volgt een vaak wat oppervlakkige beschouwing van (Amerikaanse) rechtszaken waarin sociaal-wetenschappelijke studies een rol hebben gespeeld, en de stand van het Amerikaanse recht met betrekking tot de vraag of dergelijk onderzoek toelaatbaar is als bewijs. Anders dan ten bewijze van het nut en belang van dergelijk onderzoek in abstracto geven deze voorbeelden de juridische lezer vrij weinig concrete informatie over het daadwerkelijk omgaan met - en de interpretatie van - sociaal-wetenschappelijk onderzoek.

Voor de niet-Amerikaans geïntereerde jurist en statisticus zou deze verhandeling interessant kunnen zijn, ware het niet dat de auteurs regelmatig concepten hanteren die zo specifiek Anglo-Amerikaans juridisch van aard zijn, dat zij vermoedelijk af en toe het spoor bijster zullen raken. Dit doet zich bijvoorbeeld voor in hoofdstuk 9, waarin de auteurs onder andere de ontwikkelingen die zich in Amerika hebben voorgedaan met betrekking tot de toelaatbaarheid van sociaal-wetenschappelijke studies ter ondersteuning van "motions for a change of venue". In deze studies wordt getracht de aanwezigheid of afwezigheid van vooroordelen bij de plaatselijke bevolking (de pool van potentiële juryleden) vast te stellen, teneinde daarmee vast te kunnen stellen of de beklaagde op de bedoelde locatie nog wel een eerlijk proces zal kunnen krijgen. Is dat niet het geval, dan zal de rechtszaak moeten worden verplaatst naar een andere locatie. Voor lange tijd werden de uitkomsten van dergelijke onderzoeken in Amerika niet toegelaten als bewijs, omdat rechters er vanuit gingen dat de zogenaamde "voir dire" procedure daar voldoende duidelijkheid in verschaft. De "voir dire" procedure is erop gericht potentiële juryleden naar hun oordeel over verschillende onderwerpen te vragen, ten einde te komen tot een zo objectief mogelijke jury. Daarbij merken de schrijvers op dat dergelijke onderzoeken hede ten dage toelaatbaar bewijs vormen, zolang de vragen van het onderzoek gericht is op "the interviewee's state of mind" (p. 136). In dat geval namelijk zijn de resultaten van dit onderzoek "not evidence of" "the truth of the matter asserted", and, therefore, they are not hearsay" (ibid). Om dat te begrijpen, zal de lezer bekend moeten zijn met het concept "hearsay evidence." "Hearsay" wordt in Amerika gedefinieerd als een buiten de rechtszaal afgelegde (mondelinge of schriftelijke) verklaring, zoals bijvoorbeeld een verklaring van getuige A dat hij B heeft horen zeggen dat verdachte C de roofoverval heeft gepleegd. Dergelijke verklaringen vormen in Amerikaanse rechtszaken (in principe) ontoelaatbaar bewijs. Echter, als de verklaring

niet wordt geïntroduceerd ten bewijze van het feit dat wat gezegd is daadwerkelijk is gebeurd, maar (bijvoorbeeld) om aan te tonen dat de getuige zich blijkbaar bedreigd voelde door de verdachte ("ik ben bang dat X mij zal vermoorden"), dan wordt de verklaring niet "voor waar" geïntroduceerd, en is zij wel toelaatbaar. Het is maar de vraag of niet (Amerikaanse) juristen dit soort vrij specifieke zaken zonder nadere uitleg zullen begrijpen.

De hoofdstukken 11 en 12, die beide handelen over merkenrecht, sluiten wat dat betreft veel meer aan bij de oorsprong en bedoeling van het boek. Deze hoofdstukken zijn informatiever en genuanceerder over het soort sociaal-wetenschappelijk onderzoek dat in dit soort zaken een rol kan spelen, alsook de beperkingen van dergelijk onderzoek. Daarbij wordt ook (met name in hoofdstuk 12) teruggesproken naar begrippen als standaard deviaties en P-waarden. Datzelfde geldt voor hoofdstuk 13, waarin juridische en statistische vraagstukken met betrekking tot DNA worden behandeld. Hoewel dit hoofdstuk, net als hoofdstuk 6, voor juristen niet het meest toegankelijke hoofdstuk is, is het wel een van de meest belangrijke en actuele verhandelingen voor met name strafrechtjuristen, waarbij elementaire kennis van statistiek vaak nodig gemist wordt, én regelmatig tot belangrijke misvattingen onder juristen aanleiding geeft.

Prove it with figures is uit een oogpunt van algemene ontwikkeling over het nut van sociaal-wetenschappelijk onderzoek voor juridische doeleinden niet onaardig om te lezen. Het verhaalt veel - maar vooral ook veel Amerikaans juridische - kennis. Het niveau van de (sociaal-wetenschappelijke) verhandelingen wisselt sterk tussen de verschillende hoofdstukken, hetgeen het boek wat onaantrekkelijk maakt voor het doel dat beoogd wordt. Didactisch gezien is het boek voor Continentale juristen en niet-juridisch geïnteresseerden dan ook niet het meest voor de hand liggende boek om leken te onderwijzen in sociaal-wetenschappelijk onderzoek voor juridische doeleinden.

*mr P. van Kampen, mr H. Nijboer
Vakgroep Algemeen Strafrecht
Rijksuniversiteit Leiden*

Kotz, S., Johnson, N.L.

Breakthroughs in statistics

1997, Springer-Verlag, Berlin,

590 blz, DM 114.00, ISBN 0-387-94989-5.

This third volume of the "Breakthroughs in Statistics" series contains 12 papers (out of 21) from the era 1970-1990. In this sense, Volume III focuses on more recent breakthroughs than Vols I and II (published jointly in 1992) which have only 4 (out of 39) references from that period. Volume III is organized in a similar manner as the previous volumes, namely as follows.

Each "breakthrough" is accompanied by an introductory paper written by an expert in that specific field. In general, these introductory papers are very informative and readable, even if the reader is not so familiar with the subject. An explanation is given

why the reprinted paper is considered such an important development, background information on the author is provided as are some additional references. After this introduction, the original paper is reprinted in its entirety.

Volume III includes a wide variety of papers. To give an idea, some of the discussed papers are: Yates and Cochran [1938: The Analysis of Groups of Experiments], Whittle [1953: The Analysis of Multiple Stationary Time Series], Hajek [1970: A Characterization of Limiting Distributions of Regular Estimates], and Rousseeuw [1984: Least Median of Squares Regression]. Unfortunately, two areas receive relatively much attention whereas breakthroughs in other areas are ignored. These two areas are biostatistics/epidemiology (3 papers) and Markov Chain Monte Carlo (3 papers). Notwithstanding the fact that each of these six papers is probably a breakthrough, it is somewhat disappointing that developments in other areas, such as in social statistics and industrial statistics, have not received any attention in this volume.

This has much to do with the time period considered. The more recent a paper, the harder it is to predict its long-term effect, so that there is probably less agreement among statisticians as to whether a paper is a breakthrough or not. Hence, it would be interesting to know how the editors arrived at this sample of 21 papers, but this is not clear.

This brings me to my main point of criticism, namely that the editors have failed to preserve some sort of structure in this volume (see also Ziegel, 1998). It is quite probable that within the next few years, more volumes will appear in this series with papers from the same time period as volume III. Hence, anticipating this, the editors could have categorized the recent breakthrough papers, for instance by time period or subject area. It is my opinion that, by not having done so, volume III is much like part III of a movie: it is usually not as good as the first version.

Verwijzingen:

- [1] Kotz, S., and Johnson, N.L. (1992), *Breakthroughs in Statistics Vols I and II*, Springer-Verlag: New York.
- [2] Ziegel, E. (1998), Editor's Report on: *Breakthroughs in Statistics Vol. III*, Kotz, S., Johnson, N.L. (Eds.), *Technometrics* 40, p. 165.

*Drs A.F. Huele
KdV Institute for Mathematics
University of Amsterdam*

Buldygin, V., Solntsev, S.

Asymptotic behaviour of linearly transformed sums of random variables

1997, Mathematics and its applications,
Kluwer, Dordrecht,
516 blz, fl 395.00, ISBN 0-7923-4632-7.

Readership: The book is aimed at experts in probability theory, theory of random processes with an interest in almost sure asymptotic behaviour in summability schemes.

Limit theorems for random sequences may be divided into two large parts. One of them deals with convergence of distributions (weak limit theorems). The other part deals with almost sure convergence (strong limit theorems).

This book deals with the almost sure behaviour of linearly transformed sequences of independent random variables, vectors and elements taking values in topological vector spaces.

In Part I the authors deal with series of independent random elements in the context of separable F-spaces (complete metrizable topological vector spaces which are in dual with their dual spaces) which play an important role in the whole book. Chapter 1 of the book is devoted to series of independent random elements in separable F-spaces. The main section 1.3 is devoted to the equivalence of strong and weak almost sure behaviour of series of independent symmetric summands. The authors also study Fourier analysis and convergence of series in Hilbert spaces and random series in the space of continuous functions.

The main topic of the second part of the book is concerned with the strong limit theorems for operator-normed (matrix-normed) sums of independent (and sometimes identically distributed) random vectors in finite-dimensional Euclidean spaces (Chapters 3 and 4). In Chapters 5 and 6 follow their applications to one-dimensional and multi-dimensional Gaussian Markov Processes. In Chapter 7 the authors study asymptotic properties of sample paths of random recurrent sequences. They obtain a series of interesting inequalities and limit theorems for a large number of stochastic recurrence relations of special interest.

The final Chapter 8 basically deals with the interplay between the central limit theorem and the law of the iterated logarithm in the framework of non-identically distributed summands.

The book contains a rich amount of new material that can be of interest to many researchers. Numerous sections will be of use to researchers who work with Gaussian processes, stochastic recurrence relations, operator normed sums and weighted sums. Many topics appear in this monographic form for the first time.

The exposition of the book is consistent and is a.s. self-contained. Most of the results are motivated well and proved in an elegant way and or the reader gets one or more references. The excellent bibliographic list contains over 200 titles!

I strongly recommend this book to anyone interested in strong limit theorems.

prof. dr E. Omeý
EHSAL

Lange, K.

Mathematical and statistical methods for genetic analysis

1997, Statistics for biology and health,

Springer-Verlag, Berlin,

290 blz, DM 84.00, ISBN 0-387-94909-7.

As stated in the preface, this book is written for students who are sophisticated in mathematics and statistics. Its aim is to equip the such students to engage in statistical genetics. The mathematical level of this book is high, especially in later chapters. The complexity and the breadth of the subject matter is tremendous. A variety of mathematical and statistical methods are applied to many, disparate areas of genetics. Lange's presentation is clear, but, because of his conciseness, at times demanding. This conciseness and breadth of the subject matter place a heavy burden on the reader. One could choose almost any chapter and spend a lot of time before coming to grips with the problem presented, its background, and Lange's presentation of the relevant statistical modeling.

No previous knowledge of genetics is assumed, although some basic knowledge will, I think, help. Lange devotes a chapter and the appendix to genetics proper. In view of their short length, both these chapters are surprisingly effective. Almost all genetical terms used throughout the text are either explained on the spot, or can be found, via the exhaustive index, in chapter 1, or in the appendix.

Each chapter ends with a number of problems and a list of references. The problems provided are generally quite demanding. Details that feature in these problems sometimes reappear in the main text. Solutions to (selected) problems are not provided. The list of references may serve to direct the reader to 'further reading'.

This is an impressive book, which will be of immense interest to anybody who is seriously interested in statistical modeling in genetics. Especially in view of the diversity of the subject matter, I would have preferred a less concise text. I imagine that this book could have easily been twice as long. I conclude this review with a brief summary of each chapter.

The first chapter and the appendix of this book are devoted to genetics proper. Chapter 1 is an introduction to the basic principles of population genetics. The appendix provides a resume of molecular genetics. It serves mainly to expose the reader to important concepts that are used in the rest of the book. As mentioned, both this appendix and chapter 1 are very effective.

Chapters 2 is devoted to the EM algorithm. A proof of the ascent property of the EM algorithm is provided. Following this proof, the EM algorithm is applied to allele frequency estimation and classical segregation analysis. Allele frequency estimation provides a good illustration of the EM algorithm, in which the distinction between observed and missing data, and the distinction between the E and the M step are clear.

Chapter 3 is devoted to the other major method of optimization in maximum likelihood estimation, viz. Newton's method. Related methods, the method of scoring and quasi-Newton algorithms are also discussed. The method of scoring is illustrated

in the special case of the multinomial distribution. Two examples are included: one involves the estimation of allele frequencies in the presence of inbreeding, and the other involves the calculation of expected information in linkage analysis. The chapter ends with an explanation of the Dirichlet distribution and its role in empirical Bayes' estimation of allele frequencies.

In Chapter 4 concerns multinomial problems and tests in genetics. These include tests to detect a single category with an excess number of counts, and clustering in a small number of categories. The chapter concludes with a good explanation of the transmission/disequilibrium test, which is used to test for association between a locus and a binary phenotype. The rationale behind these tests is explained well.

Chapter 5 treats kinship and genetic identity coefficients, which are used to determine the degree of genetic relatedness among individuals. In chapter 6 these are used in various applications including genotype prediction (predicting the genotype of person i from one or more of his relatives), deriving the covariances for a quantitative trait, calculating genetic risk ratio's (probability of person i being affected, given an affected relative), and linkage analysis.

Chapter 7 is devoted to the computation of Mendelian Likelihoods. This involves testing Mendelian hypothesis in pedigrees. Mendelian hypotheses, relating to a finite number of major gene, are in themselves simple, the computational difficulties in obtaining maximum likelihood estimates can be considerable. Various strategies to facilitate computation are presented. A number of interesting examples are provided, including paternity testing.

Chapter 8 concerns the polygenic model, i.e., a model for phenotypes that are influenced by many genes. The method is presented for a (multivariate) normal phenotype. The loglikelihood function, its derivatives, and the information matrix are presented. The polygenic model is extended to accommodate dichotomous phenotypes (the polygenic threshold model) and polygenic influence in the presence of major gene. This gives rise to a mixture of normals, in which the components correspond to the genotypes of the major gene.

Markov Chain Monte Carlo (MCMC) Methods are explained in chapter 9. Discrete time markov chains, the metropolis algorithm, and simulated annealing are discussed. Descent graphs are introduced to model the transmission of alleles through a pedigree. These are treated as states in applying MCMC to the calculation of locations scores in linkage analysis and haplotyping.

Chapter 10 is devoted to the reconstruction of evolutionary trees. The objective here is to determine the evolutionary relationship between taxa, i.e. how (and how fast) they diverged from common ancestors. The raw data used in such reconstruction are DNA sequences. A maximum likelihood reconstruction method is developed and used in determining the origin of eukaryotes.

Chapter 11 deals with radiation hybrid mapping, a technique designed to determine the order of loci and the distances between adjacent loci. A hybrid rodent cell is created including a single human chromosome. Radiation breaks up the chromosome, which are subsequently repaired by the cell repair mechanisms. During repair the fragments

of human chromosome are inserted into the rodent's chromosomes. The presence of two loci on a single fragment suggest that the loci are closely linked. Two types of hybrid cells are discussed: haploid radiation hybrids and polyploid radiation hybrids. The minimum obligate breaks criterion is used to identify the order of loci. Lange presents maximum likelihood methods which are more informative in that they provide physical maps and a comparison between putative orders. The chapter ends with the presentation of a Bayesian method.

Chapter 12 concerns models of recombination. The distance between loci on a chromosome are expressed by means of mapping functions which relate the distance to the probability of recombination between the loci. Models of recombinations are presented based on the poisson distribution and on renewal models. Recombination events are known not to be independent. A model for recombination is presented that takes this into account.

The final chapter introduces Poisson approximations based on the so-called Chen-Stein method for approximating the distribution of a sum of weakly dependent Bernoulli random variables. Areas of application include the determination of the largest gap between to markers after randomly placing markers throughout the human genome, DNA sequence matching (do two DNA sequences share significant similarities?).

C. Dolan

Vakgroep Psychonomie

Vrije Universiteit Amsterdam

Lindsey, J.K.

Applying generalized linear models

1997, Springer texts in statistics,

Springer-Verlag, Berlin,

280 blz, DM 98.00, ISBN 0-387-98218-3.

Om gelijk met de deur in huis te vallen: ik vind het compact geschreven boek "Applying Generalized Linear Models" van J.K. Lindsey een welkome aanvulling op de bestaande literatuur op het gebied van gegeneralizeerde lineaire modellen (g.l.m.'s). Het benadrukt een modelmatige aanpak, "without becoming lost in problems of statistical inference". Hier ligt ook een bezwaar: de toegepast statisticus die de beschreven modellen wil toepassen, krijgt weinig suggesties hoe hij of zij dat precies moet doen. In die zin is ook de titel misleidend: het boek is niet zozeer een gids voor het toepassen van g.l.m.'s met uitvoer van een softwarepakket en al, maar een compilatie van modellen, waaraan de lezer kan ruiken. De auteur geeft in de inleiding toe, dat de lezer via dit boek niet een expert zal worden, maar hopelijk inzicht zal krijgen in de achterliggende eenheid. Om een expert te worden zal andere literatuur bestudeerd moeten worden. Gelukkig worden aan het eind van ieder hoofdstuk referenties gegeven naar recente literatuur.

Voor wie is het boek interessant? Iedere toegepast statisticus, die met enige regelmaat g.l.m.'s gebruikt (en wie doet dat niet?), zou dit boek kunnen waarderen. Als tekst

voor een cursus statistisch modelleren lijkt het me nogal hoog gegrepen. De auteur ziet hier wel mogelijkheden, met name vanwege de opgaven, waarin vele datasets worden gepresenteerd. De overvloed aan data uit velerlei vakgebied is een aantrekkelijk punt van het boek.

Dan nu een korte beschrijving van de inhoud. Hoofdstuk 1 beschrijft g.l.m.'s in algemene zin. Om het niveau aan te geven: het hoofdstuk eindigt met een opsomming van uitbreidingen van g.l.m.'s, zoals geparametriseerde linkfunctie, variantiefunctie, samengestelde linkfuncties, niet-lineaire structuur, m.i. onderwerpen waar de meest gangbare boeken over g.l.m.'s mee zullen eindigen. Voor het vergelijken van (ook niet-geneste) modellen wordt Akaike's Informatie Criterium (AIC) gebruikt. Hoofdstuk 2 behandelt discrete data, met als belangrijkste onderwerpen de samenhang tussen log lineaire en logistische modellen, "models of change", overdispersie, random effects en (leuk!) het Rasch model. Hoofdstuk 3 vond ik erg verrassend: er wordt beschreven hoe met Poisson regressie modellen een groot aantal verdelingen gefit kunnen worden, inclusief afgeknotte en multivariate verdelingen. Hoofdstuk 4 beschrijft het fitten van exponentiële, logistische en Gompertz-groeikrommen zonder rekening te houden met seriële correlatie. Een voorbeeld van een complexer model betreffende AIDS wordt uitgewerkt. Hoofdstuk 5 stipt discrete tijdreeksen aan: Poisson- en Markov-processen en repeated measurements. Hoofdstuk 6 beschrijft in enkele pagina's de analyse van survival data: niet-parametrisch, parametrisch en semi-parametrisch. De auteur houdt het kort omdat de modellen niet recht-toe-recht-aan g.l.m.'s zijn. Hoofdstuk 7 behandelt de analyse van reeksen gebeurtenissen (event histories, telprocessen). Hoofdstuk 8 beschrijft spatieel modelleren als een onderwerp waarvoor g.l.m.'s niet goed geschikt zijn. Hoofdstuk 9 behandelt lineaire regressie (response surface), ANOVA en niet-lineaire regressie, niet direct een onderwerp dat men achterin dit boek zou verwachten. Het moeilijkste onderwerp is tot het laatst bewaard: in hoofdstuk 10 komen dynamische modellen aan bod. Appendix A behandelt bondig likelihood inferentie, inclusief frequentistische en Bayesiaanse "decision-making". In appendix B komen diagnostische technieken aan de orde.

*G. Gort
Vakgroep Wiskunde
Landbouw Universiteit Wageningen*

Klein, J.P., Moeschberger, M.L.

Survival analysis. Techniques for censored and truncated data

1997, Statistics for biology and health,

Springer-Verlag, Berlin,

xiv+502 blz, DM 94.00, ISBN 0-387-94829-5.

Dit boek gaat over de analyse van wachttijden op een bepaalde gebeurtenis, in context van medische toepassingen. De waarnemingen kunnen gecensureerd en afgekapt (truncated) zijn, hetgeen om een speciale aanpak vraagt. Van censurering is er sprake

als men, in plaats van de exacte wachttijd, een onder- of bovengrens op de wachttijd waarneemt. Afkapping houdt in dat de wachttijd alleen binnen bepaalde grenzen waargenomen kan worden.

De stof is gepresenteerd in 13 hoofdstukken. De eerste drie introduceren de problematiek en basisbegrippen. Het boek begint met een bespreking van 18 datasets die verder in het boek als voorbeelden gebruikt worden. De meeste van deze datasets staan in het boek, de lezer kan ze echter ook van de Web site van de uitgever halen. Alle voorbeelden, op n na ("Time to First Use of Marijuana"), hebben een medische context. Hoofdstuk 2 bespreekt de basale kansmodellen en parameters. In Hoofdstuk 3 komen de vormen van censurering aan bod. Dit hoofdstuk eindigt met 9 bladzijden over telprocessen. De lezer kan deze helder gepresenteerde maar wel ingewikkelde theorie desgewenst overslaan.

Hoofdstukken 4, 5 en 6 gaan over het schatten van overlevingsfunctie en gerelateerde grootheden. Aan bod komen onder andere betrouwbaarheidsbanden, het schatten van 'excess mortality' ten opzichte van een standaard, gladde schatters en Bayesiaanse methoden. Hoofdstuk 7 is gewijd aan het toetsen in het geval van n , twee of meerdere steekproeven. Naast de logranktoets, de gegeneraliseerde toets van Wilcoxon en hun trend- en gestratificeerde varianten bespreekt het boek ook de toets van Renyi en allerlei minder bekende methoden.

Hoofdstukken 8 tot en met 12 gaan over regressiemodellen. Hoofdstukken 8 en 9 over het model van Cox. Aandacht wordt besteed aan onder andere gevolgen van knopen voor de partiele aannemelijkheidsfunctie, aan model-building en aan tijdsafhankelijke verklarende variabelen. Hoofdstuk 10 behandelt het model van additieve hazardfuncties, Hoofdstuk 11, heel uitvoerig, de regressie-diagnostiek en Hoofdstuk 12 de parametrische regressie-modellen. In het laatste hoofdstuk komen multivariate wachttijden aan bod, onder andere het gamma frailty model.

Aan het eind van de meeste paragrafen vindt men aanvullende informatie. Diepere theoretische achtergronden, zoals relatie tot de theorie van telprocessen, in Theoretical Notes, en praktische aanwijzingen, over verwante methoden of over de programmatuur, in Practical Notes. Elk hoofdstuk eindigt met op de praktijk gerichte opgaven. Het boek bevat verder aanhangsels over numerieke en theoretische aspecten van de maximum likelihood methode, statistische tabellen, en tabellen met datasets.

De presentatie van de stof is degelijk en overzichtelijk, met veel uitgewerkte voorbeelden en met goed bruikbare opgaven. Dit maakt het geschikt voor onderwijs. Het veronderstelt wel een goede basiskennis van statistiek. Dankzij zijn volledigheid is het boek ook geschikt als naslagwerk voor statistici. Samenvattend: voor een ieder die serieus met overlevingsduren aan de slag wil is dit boek een aanrader.

*dr V. Fidler
Vakgroep Medische Statistiek
Rijksuniversiteit Groningen*

Gerber, S., Voelkl, K.E.

The SPSS guide to The new statistical analysis of data

1997, Springer-Verlag, Berlin,

200 blz, DM 48.00, ISBN 0-387-94821-X.

Dit boek is geschreven als een gids bij het boek *The New Statistical Analysis of Data* van T.W. Anderson en J.D. Finn, maar kan, volgens de auteurs, ook zelfstandig als inleiding in SPSS gebruikt worden. De gids houdt de volgorde van de onderwerpen van het boek aan en gebruikt dezelfde voorbeelden. Het is gebaseerd op de versie 6.1 van SPSS en dus enigszins gedateerd. In nieuwere versies zijn sommige vensters veranderd, zijn grafische mogelijkheden verbeterd en is de uitvoer beter georganiseerd. Maar voor wat betreft de basale statistiek zijn er geen inhoudelijke verschillen.

De stof is gepresenteerd in 16 hoofdstukken, die gegroepeerd zijn in vijf delen. Het eerste deel behandelt de basale handelingen, zoals invoeren van data, manipulatie met files, transformeren van variabelen, het printen. De overige vier delen zijn Descriptive Statistics (63 blz), Probability (17 blz), Statistical Inference (47 blz) en Statistical Methods for Other Problems (45 blz). Elk hoofdstuk eindigt met eenvoudige opgaven en een overzicht van SPSS-syntax. De gebruikte data zijn op te halen op de Web site van de uitgever. Het boek is goed verzorgd. Ietwat merkwaardig vind ik dat in het boek tabellen als figuren worden aangeduid.

De gids bestrijkt de traditionele inleiding tot statistiek, tot en met enkelvoudige lineaire regressie en variantie analyse. Ik vermoed dat ook het boek van Anderson en Finn nogal klassiek van structuur is, ondanks zijn titel. Zo een structuur leent zich slecht voor een op praktijk gerichte introductie van een pakket als SPSS. Een voorbeeld. Het boek van Anderson en Finn wijdt kennelijk aparte paragrafen aan kengetallen zoals de modus, het gemiddelde, de mediaan, de spreidingsbreedte, de gemiddelde absolute afwijking (!) en de standaardafwijking. De gids volgt dit op de voet, en zo heeft elk kengetal een eigen paragraaf, compleet met gedetailleerde instructie hoe het aan het programma te ontfutselen is. Deze werkwijze leidt tot onnodige herhaling. Aan de interpretatie van de kengetallen wordt daarbij geen aandacht besteed. De kennelijk ongewenste maar wel door SPSS geleverde kengetallen, bijvoorbeeld de kurtosis, worden niet eens genoemd. De auteurs wijzen wel op eventuele tekortkomingen van SPSS, bijvoorbeeld op de soms foutieve berekening van kwartielen.

In het deel over kansrekening verricht de lezer aselechte trekkingen uit enkele kansverdelingen en wordt de Centrale limietstelling geadstrueerd. Voor dat laatste is een speciale file (van de Web site) nodig. De gids geeft helaas geen andere instructie voor het aanmaken van een file met een gekozen aantal cases dan deze één voor één in te typen. Dit beperkt aanzienlijk, en onnodig, de simulatie mogelijkheden.

In het deel over verklarende statistiek vindt men onder andere de toets voor het vergelijken van gemiddelden in het geval dat de standaardafwijking bekend is. Omdat SPSS dit niet kan, voert men de berekening met de hand. Ik betwijfel of dit een aanvulling op het boek van Anderson en Finn is. Zoals reeds boven gesteld leidt het te getrouw volgen van dit boek tot kunstmatige situaties. Overigens is er in de praktijk

een aanvullende berekening dikwijls nodig en zinvol, en het is op zich positief dat de gids laat zien dat er met de uitvoer van SPSS ook nog gerekend moet worden.

Samenvattend: deze gids kan misschien enige nut hebben als aanvulling op het boek van Anderson en Finn, maar is niet geschikt om zelfstandig gebruikt te worden.

dr V. Fidler

*Vakgroep Medische Statistiek
Rijksuniversiteit Groningen*

Wilcox, R.R.

Introduction to robust estimation and hypothesis testing

1997, Statistical modelling and decision science,

Academic Press, San Diego, California,

xii+296 blz, \$ 59.95, ISBN 0-12-751545-3.

Na het pionierswerk in de zestiger jaren van Huber en Hampel en het vele onderzoek dat daar in de afgelopen decennia op volgde, zijn robuuste methoden momenteel alom geaccepteerd in de wetenschappelijke onderzoekswereld. Voor toegepaste onderzoekers kan het desalniettemin lastig zijn robuuste methoden toe te passen, omdat de relevante software nauwelijks in standaard software pakketten aanwezig is en omdat de tekstboeken op dit gebied moeilijk leesbaar zijn.

Dit boek biedt uitkomst. Het doel ervan is volgens de auteur om een relatief eenvoudige beschrijving te geven van beschikbare methoden die op grond van theoretische en simulatie resultaten als de beste worden beschouwd. Er wordt helder uitgelegd waarom robuuste methoden belangrijk zijn in de statistiek. Roubuuste alternatieven voor een aantal klassieke schatters en toetsen worden beschreven. Voor wie de beschikking heeft over S-PLUS(is dit boek helemaal een feest. Op de website van Academic Press is namelijk een bibliotheek te vinden met alle S-PLUS functies die in dit boek zijn beschreven.

De inhoud van het boek is globaal als volgt. In Hoofdstuk 1 staan de belangrijkste praktische argumenten waarom robuuste methoden nodig zijn. Hoofdstuk 2 is het meest technisch van het hele boek. Hierin wordt het theoretisch fundament van robuuste methoden beschreven. Aan de orde komen onder andere de concepten: qualitative robustness, infinitesimal robustness, quantitative robustness, influence function, breakdown point, M-estimators. In hoofdstuk 3 worden robuuste locatie schatters en spreidingsmaten besproken, hoofdstuk 4 gaat over betrouwbaarheidsintervallen behorend bij de schatters van het voorgaande hoofdstuk en hoofdstuk 5 over het vergelijken van twee groepen. In Hoofdstuk 6 worden enkele veel voorkomende ANOVA designs behandeld en hoofdstuk 7 gaat over het schatten van correlaties en het opsporen van multivariate uitbijters. De laatste twee hoofdstukken behandelen regressie analyse. Naast de (gegeneraliseerde) M-estimators worden high-breakdown schatters besproken en tenslotte nog enkele alternatieve schattingsprocedures. In Hoofdstuk 9 wordt ingegaan op heteroscedastische storingstermen en het vinden van niet-lineaire verbanden. Ieder hoofdstuk wordt afgesloten met een aantal oefeningen. Het boek maakt

duidelijk dat niet één methode in alle situaties de beste is; dit is afhankelijk van welke criteria worden gehanteerd. Zeker is dat de klassieke methoden op veel criteria slecht scoren.

Vanaf hoofdstuk 3 wordt iedere inhoudelijke paragraaf, waarin een schatter of een toets wordt beschreven, gevolgd door een paragraaf waarin n of meer S-PLUS functies uitgebreid worden beschreven en meestal met een voorbeeld worden geïllustreerd. S-PLUS bezitters kunnen dus onmiddellijk aan de slag. Een nadeel is dat wie niet over S-PLUS beschikt en onbekend is met de structuur ervan weinig heeft aan deze paragrafen, die bij elkaar een aanzienlijk deel van het boek beslaan.

Het boek is vooral erg praktisch van opzet en in het algemeen goed leesbaar. Er zijn veel voorbeelden en veel recepten hoe schatters en toetsen moeten worden uitgerekend. Deze recepten staan in "How to compute"-tabellen waarin een berekeningsmethode volledig wordt vermeld. Naast deze praktische kant geeft het boek toch ook een aardig theoretisch beeld van wat het begrip robuust in de statistiek betekent. Mensen die zich met het analyseren van data bezig houden, zouden veel aan dit boek kunnen hebben. Ook voor mensen die kennis willen maken met de robuuste statistiek lijkt dit boek geschikt. Het boek of delen van het boek zouden ook zeker niet misstaan in een onderwijs-module voor studenten die later te maken kunnen krijgen met data analyse (bijvoorbeeld in de sociale wetenschappen). Enige basiskennis van de statistiek is vereist, in het bijzonder kennis over hypothesen toetsen, ANOVA, en regressie analyse.

P. Verboon

Sector Statistische Methoden

Centraal Bureau voor de Statistiek

Binnengekomen boeken

Hieronder volgt een overzicht van de recente uitgaven die sinds de vorige aflevering van *Kwantitatieve Methoden* bij de redactie zijn binnengekomen. Mocht er belangstelling bestaan voor het bespreken van een boek, dan kan contact opgenomen worden met de boekbesprekingsredacteur. Indien het boek dan nog niet vergeven is, krijgt de kandidaat het boek gratis thuisgezonden. Hiervoor dient dan binnen een half jaar een recensie te worden teruggezonden. Een aantal titels is niet fysiek toegezonden, maar kan door ons besteld worden bij de uitgever. Dit heeft consequenties voor de levertijd van deze titels.

Suresh, N.C., Kay, J.M. (eds)
Group technology and cellular manufacturing
 1997, Kluwer, Dordrecht,
 568 blz, fl 370.00, ISBN 0-7923-8080-0.

Mistakidis, E.S., Stavroulakis, G.E.
Algorithms, heuristics and engineering applications by the F.E.M.
 1997, Nonconvex optimization and its applications,
 Kluwer, Dordrecht,
 300 blz, fl 260.00, ISBN 0-7923-4812-5.

Kreinovich, V., Lakeyev, A., Rohn, J., Kahl, P.
Computational complexity and feasibility of data processing and interval computations
 1997, Applied optimization,
 Kluwer, Dordrecht,
 472 blz, fl 375.00, ISBN 0-7923-4865-6.

Abramov, A.P.
Connectedness and necessary conditios for an extremum
 1998, Mathematics and its applications,
 Kluwer, Dordrecht,
 212 blz, fl 175.00, ISBN 0-7923-4910-5.

Kaluzny, S., Vega, S.C., Cardoso, T.P., Shelly, A.A.
S+SpatialStats

1997, Springer-Verlag, Berlin,
 385 blz, DM 88.00, ISBN 0-387-98226-4.

Via de VVS
 25% korting!

Cobb, G.W.

Introduction to design and analysis of experiments

1998, Springer-Verlag, Berlin,
 795 blz, DM 114.00, ISBN 0-387-94607-1.

Via de VVS
 25% korting!

Foster, D.P., Stine, R.A., Waterman, R.P.

Basic business statistics

1998, Springer-Verlag, Berlin,
 xvi+244 blz, DM 69.00, ISBN 0-387-98354-6.

Via de VVS
 25% korting!

Foster, D.P., Waterman, R.A., Waterman, R.P.

Business analysis using regression

1998, Springer-Verlag, Berlin,
 348 blz, DM 79.00, ISBN 0-387-98356-2.

Via de VVS
 25% korting!

Anderson, D.R., Sweeney, D.J., Williams, T.A.

Statistiek voor economie en bedrijfskunde

1998, Academic Service, Schoonhoven,
 507 blz, fl 75.00, ISBN 90-5261-250-1.

Cramer, J.S., Hemelrijk, J.

Statistiek eenvoudig

1998, Academic Service, Schoonhoven,
 127 blz, fl 19.90, ISBN 90-5712-011-9.

Draper, N.R., Smith, H.

Applied regression analysis 3e

1998, Wiley series in probability and statistics,
 John Wiley & Sons Ltd., Chichester,
 xvii+706 blz, £ 45.00, ISBN 0-471-17082-8.

Johnson, D.

Applied multivariate methods for data analysis

1998, Duxbury Press, Belmont, California,
425 blz, £ 25.99, ISBN 0-534-23796-7.

Berk, K., Carey, P.

Data analysis with Microsoft Excel. Windows 95 edition

1998, Duxbury Press, Belmont, California,
512 blz, £ 21.00, ISBN 0-534-52929-1.

Kleinbaum, D.G., Kupper, L.L., Muller, K.E.

Applied regression analysis and multivariable methods

1998, Duxbury Press, Belmont, California,
736 blz, £ 25.99, ISBN 0-534-20910-6.

Wedel, M., Kamakura, W.A.

Market segmentation: conceptual and methodological foundations

1998, International series in quantitative marketing,
Kluwer, Dordrecht,
378 blz, ISBN 0-7923-8071-1.

Dembo, A., Zeitouni, O.

Large deviations techniques and applications

1998, Applications of mathematics,
Springer-Verlag, Berlin,
420 blz, DM 124.00, ISBN 0-387-98406-2.

Via de VVS
25% korting!

Rachev, S.T., Rüschendorf, L.

Mass transportation problems. Part I: theory.

1998, Probability and its applications,
Springer-Verlag, Berlin,
520 blz, DM 158.00, ISBN 0-387-98350-3.

Via de VVS
25% korting!

Rachev, S.T., Rüschendorf, L.

Mass transportation problems. Part II: applications.

1998, Probability and its applications,

Springer-Verlag, Berlin,

520 blz, DM 158.00, ISBN 0-387-98352-X.

Via de VVS

25% korting!

Rossmann, A.J., Chance, B.L.

Workshop statistics: discovery with data and Minitab

1998, Springer-Verlag, Berlin,

xxviii+486 blz, DM 79.00, ISBN 0-387-98411-9.

Via de VVS

25% korting!

Rawlings, J.O., Pantula, S.G., Dickey, D.A.

Applied regression analysis. A research tool

1998, Springer texts in statistics,

Springer-Verlag, Berlin,

670 blz, DM 168.00, ISBN 0-387-98454-2.

Via de VVS

25% korting!

Schmidt, W.H., Heier, K., Bittner, L., Bulirsch, R. (eds)

Variational calculus, optimal control and applications

1998, International series of numerical mathematics,

Birkhäuser, Basel,

342 blz, sFr 138.00, ISBN 3-7643-5906-4.

Desch, W., Kappel, F., Kunisch, K. (eds)

Control and estimation of distributed parameter systems

1998, International series of numerical mathematics,

Birkhäuser, Basel,

310 blz, sFr 128.00, ISBN 3-7643-5835-1.

Vorotnikov, V.I.

Partial stability and control

1998, Birkhäuser, Basel,

432 blz, sFr 158.00, ISBN 3-7643-3917-9.

Eberlein, E., Hahn, M., Talagrand, M. (eds)

High dimensional probability

1998, Progress in probability,
Birkhäuser, Basel,
344 blz, sFr 128.00, ISBN 3-7643-5867-X.

Abraham, B. (ed)

Quality improvement through statistical methods

1998, Statistics for industry and technology,
Birkhäuser, Basel,
440 blz, sFr 138, ISBN 3-7643-4052-5.

Cieslik, D.

Steiner minimal trees

1998, Nonconvex optimization and its applications,
Kluwer, Dordrecht,
324 blz, fl 260.00, ISBN 0-7923-4983-0.

Gang Yu (ed)

Operations research in the airline industry

1997, International series in operations research and management science,
Kluwer, Dordrecht,
496 blz, fl 370.00, ISBN 0-7923-8039-8.

Shor, N.Z.

Nondifferentiable optimization and polynomial problems

1998, Nonconvex optimization and its applications,
Kluwer, Dordrecht,
412 blz, fl 330.00, ISBN 0-7923-4997-0.

van Maanen, H.

FC Algebra. Cijfers en sport

1998, Boom, Meppel,
200 blz, fl 29.50, ISBN 90-5352-4088.

Waller, G., Meehl, P.E.

Multivariate taxometric procedures. Distinguishing types from continua

1998, Advanced quantitative techniques in the social sciences series,

Sage, Thousand Oaks, California,

x+150 blz, £ 21.00, ISBN 0-7619-0257-0.

Fennema, F., van der Eijk, C., Schijf, H.

In search of structure. Essays in social science and methodology

1998, Het Spinhuis, Amsterdam,

277 blz, fl 65.00, ISBN 90-5589-114-2.

Bickel, P.J., Klaassen, C.A.J., Ritov, Y., Wellner, J.A.

Efficient and adaptive estimation for semiparametric models

1998, Springer texts in statistics,

Springer-Verlag, Berlin,

585 blz, DM 104.00, ISBN 0-387-98473-9.

Via de VVS

25% korting!

Anderson, D.R., Sweeney, D.J., Williams, T.A.

Opgaven bij Statistiek voor economie en bedrijfskunde

1998, Academic Service, Schoonhoven,

154 blz, ISBN 90-5261-251-X.

Barros, A.I.

Discrete and fractional programming techniques for location models

1998, Combinatorial optimization,

Kluwer, Dordrecht,

196 blz, fl 180.00, ISBN 0-7923-5002-2.

Giannessi, F., Komlósi, S., Rapcsák, T. (eds.)

New trends in mathematical programming. Hommage to Steven Vajda

1998, Applied optimization,

Kluwer, Dordrecht,

328 blz, fl 260.00, ISBN 0-7923-5036-7.

Babuška, R.

Fuzzy modeling for control

1998, Intelligent technologies,
Kluwer, Dordrecht,
288 blz, fl 280.00, ISBN 0-7923-8154-8.

Reemtsen, R., Rückmann, J.-J. (eds.)

Semi-infinite programming

1998, Nonconvex optimization and its applications,
Kluwer, Dordrecht,
428 blz, fl 330.00, ISBN 0-7923-5054-5.

Drexl, A., Kimms, A. (eds.)

Beyond manufacturing resource planning (MRP II). Advanced models and methods for production planning

1998, Springer-Verlag, Berlin,
419 blz, DM 148.00, ISBN 3-540-64247-1.

Via de VVS

25% korting!

Smith, P.J.

Into statistics. A guide to understanding statistical concepts in engineering and the sciences

1998, Springer-Verlag, Berlin,
550 blz, DM 98.00, ISBN 981-3083-89-1.

Via de VVS

25% korting!

Chambers, J.M.

Programming with data. A guide to the S language

1998, Springer-Verlag, Berlin,
480 blz, DM 84.00, ISBN 0-387-98503-4.

Via de VVS

25% korting!

Golden, K.M., Grimmett, G.R., James, R.D., Milton, G.W.

Mathematics of multiscale materials

1998, The IMA volumes in mathematics and its applications.,
Springer-Verlag, Berlin,
300 blz, DM 129.00, ISBN 0-387-98528-X.

Via de VVS

25% korting!

De door Springer gepubliceerde boeken in bovenstaande lijst kunnen door individuele leden met een korting van 25% op de officiële prijs¹ aangeschaft worden via de Centrale Administratie van de VVS. Bestelformulieren zijn te verkrijgen via Internet, URL <http://www.vvs-or.nl/boeken/>.

¹De prijzen zoals vermeld in de lijst van binnengekomen boeken kunnen foutief zijn. Een betrouwbaarder bron van informatie is de Springer website <http://www.springer.de/>.

Oproep voor boekbesprekers

De redactie van *Kwantitatieve Methoden* is steeds op zoek naar mensen die bereid zijn een pas verschenen boek te lezen met het doel daarover een oordeel te geven. Zo'n recensie verschaft de leden van de VVS en overige lezers van *Kwantitatieve Methoden* inzicht in de kwaliteiten van het boek en verschaft nuttige informatie bij een overwogen aanschaf. Voor de schrijvers biedt een kritisch oordeel van het publiek waarvoor het boek geschreven een mogelijkheid tot verdere verbetering. De uitgever heeft uiteraard meer baat bij een lovende recensie. Als tegenprestatie mogen de recensenten de besproken boeken behouden.

Behalve de algemene boeken over statistiek of besliskunde, waar veel mensen zinnige dingen over kunnen zeggen, behandelen veel boeken een specialistisch onderwerp. Als wij, soms op de gok, iemand een dergelijk boek toesturen, blijkt een enkele keer dat 'het boek precies aansluit bij het onderzoek' van de recensent. Dat is prettig, omdat veel mensen daarbij voordeel hebben. De recensent, omdat hij/zij het nieuw verschenen boek op het onderzoeksgebied ter beschikking heeft, en de auteur(s) omdat er een terzake kundig oordeel wordt gegeven. Om nu de succeskans van het toesturen van een boek te vergroten, willen wij graag in contact komen met zoveel mogelijk collega statistici, kansrekenaars en besliskundigen. Stuur het onderstaande antwoordstrookje naar het adres van de boekbesprekingsrubriek als u interesse heeft om in de toekomst een boek te bespreken. Als er dan in de toekomst wellicht een gloednieuw boek op uw onderzoeksgebied binnenkomt hebben wij een grotere kans op een perfecte match.

Ik ben graag bereid een boek op mijn interessegebied te bespreken.

Naam	
Adres	
Telefoon	
e-mail	
Interessegebieden (graag nauwkeurig omschrijven)	