

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. 65

Bapat, R.B.

Linear algebra and linear models

Haight, J.

Taking chances. Winning with probability

Vapnik, V.N.

The nature of statistical learning theory

Wilkinson L.

The grammar of graphics

Ross, S.M.

Introduction to the mathematics of finance: options and other topics

Beltrami, E.

What is random? Chance and order in mathematics and life

Rao, C.R., Toutenburg, H.

Linear models. Least squares and alternatives

Efromovich, S.

Nonparametric curve estimation

Serfozo, R.

Introduction to stochastic networks

Leonard, T., Hsu, J.S.J.

Bayesian methods: an analysis for statisticians and interdisciplinary researchers

Coxon, A.P.M.

Sorting data: Collection and analysis

Kallianpur, G., Karandikar, R.L.

Introduction to option pricing theory

Le Gall, J.-F.

Spatial branching processes, random snakes and partial differential equations

Padberg, M.

Linear optimization and extensions

Tayur, S., Ganeshan, R., Magazine, M. (eds)

Quantitative models for supply chain management

Hooley, G., Hussey, M.

Quantative methods in marketing

Johnson, R., Kubby, P.

Just the essentials of elementary statistics

Huygens, C.

Van rekeningh in spelen van geluck

Aven, T., Jensen, U.

Stochastic models in reliability

Gilliam, D.S., Picci, G.

Dynamical systems, control, coding, computer vision. New trends, interfaces and interplay

van den Broek, J., Kop, P.

Kattenaids en statistiek

Boertjens, J.H.S., Kroes, J.J.

Internetgids gezondheid

Akaike, H., Kitagawa (eds)

The practice of time series analysis

Belsley, D.A.

Conditioning diagnostics; collinearity and weak data in regression

Bapat, R.B.

Linear algebra and linear models

2000, Springer-Verlag, Berlin,

x+138 blz, DM 89.00, ISBN 0-387-98871-8.

Via de VVS
25% korting!

The monograph "Linear algebra and linear model" is a compact but thorough introduction to the basic concepts of linear algebra and linear models with emphasis on the theory of linear estimation and linear hypotheses testing. The theory and proofs are described using the "algebraic" approach instead of the "projections" approach, with focus on proving the results.

This edition is thoroughly revised and enlarged in contrast to the first edition of 1993. Besides, two sections, one on the "volume of a matrix" and another on "star order", as well as some exercises are added. The book contains six chapters:

- Vector Spaces and Matrices
- Linear Estimation
- Tests of Linear Hypotheses
- Singular Values and Their Application
- Block Designs and Optimality
- Rank Additivity

Each chapter begins with an outline of some theoretical aspects of matrix theory followed by a practical application. For example, in chapter 4, singular value decomposition and majorization are described first, followed by an outline of principal components. Each topic is completed with a paragraph of exercises. And it was good to discover that the author had not forgotten to give most of the solutions too.

A short overview of the basic concepts of matrix theory such as vector spaces, dimension, orthogonality, singularity, eigenvalues and the Frobenius inequality are given in chapter 1. The next chapter is devoted to the estimation of linear models, weighing designs and the role of generalized inverses. Chapter 3 discusses variance analysis, multiple correlation, regression models and the general linear hypothesis.

Chapter 4 reviews singular value decomposition, majorization, extremal representations and how to apply these concepts in the calculation of the principal components or canonical correlation. An introduction to the main concepts of design theory is given in chapter five with emphasis on the diverse optimality criterions. The last chapter reviews the properties of generalized inverses (rank additivity) and one of its applications (the General Linear Model).

Who should read this monograph? According to the author advanced undergraduate and master students with interest in the theoretical aspects of above described topics. Some background in mathematics is desirable as well, because of the amount

of "theory" and proofs. This book is less suitable for students with interest in the practical aspects of the described topics. Apart from students, it may be interesting for teachers in mathematics or statistics as a source of standard results and problems.

In summary, the monograph is well, clearly and compactly written. It contains a review of the main aspects of linear algebra and linear models and it has succeeded in connecting the two topics, completed with a lot of exercises. Besides, it can serve as a book of reference.

*M.H.J. de Bruijne
Capaciteitsgroep Medische Statistiek
Leids Universitair Medisch Centrum*

Haight, J.

Taking chances. Winning with probability

2000, Oxford University Press, Oxford,

344 blz, £ 8.99, ISBN 0-19-850291.

De auteur stelt in het voorwoord dat de meeste mensen die meespelen in een loterij, kaarten of dobbelspelletjes spelen, zich wel bewust zijn van kansen. Ook als we geen berekeningen maken, hebben we wel een idee over wat mogelijk en waarschijnlijk is. Meestal zijn deze ideeën voor praktisch gebruik wel voldoende, maar zonder een goed begrip van het combineren en interpreteren van kansen is het makkelijk om te blunderen.

Doel en doelgroep De schrijver wil het begrip van waarschijnlijkheid bij de lezer vergroten door hem bloot te stellen aan veel spelsituaties waarin toeval een rol speelt. Als er dan toch een foute beslissing wordt genomen, dan is het duidelijker waarom het fout ging. Het boek richt zich op de 'wiskundige leek', die geïnteresseerd is in statistisch redeneren. Volgens de auteur zijn er geluiden dat elke wiskundige uitdrukking in een populair boek het lezerspubliek halveert. Samen met het doel dat de auteur voor ogen heeft, zorgt ervoor dat de wiskundige symbolen in de tekst vrij beperkt in aantal zijn. Ter compensatie is er een aantal appendici waarin het een en ander meer wiskundig wordt uitgewerkt (tellen, waarschijnlijkheid, gemiddelden en spreiding, goodness-of-fit tests, kelly strategy).

Onderwerpen Het boek begint met een inleiding over het begrip waarschijnlijkheid. Aardig hierin is dat de auteur erkent dat mensen verschillende kansen toekennen aan dezelfde gebeurtenis (bijvoorbeeld tijdreizen). De begrippen elkaar uitsluitende gebeurtenissen, onafhankelijkheid, gemiddelde en variabiliteit worden uitgelegd. Daarna passeren een aantal onderwerpen de revue zoals loterijen, voetbalpools, spelletjes met munten en dobbelstenen, televisiespelletjes en casinospelen. De settings variëren van de natuur, bedrijfsleven, sport tot vrijetijd. In het hoofdstuk over dobbelstenen staat als voorbeeld het monopolyspel genoemd. Het vakje Gevangenis wordt vaker bezocht dan elk ander vakje. Bijna twee keer zoveel, bleek uit een simulatie. Samen met

de kansverdeling van de som van de ogen van de gegooide dobbelstenen, zorgt dit ervoor dat bepaalde steden vaker bezocht worden. Als dat weer gecombineerd wordt met de verwachte opbrengst min de verwachte kosten, dan blijkt dat bepaalde steden interessanter zijn om te bezitten dan andere. In tegenstelling tot wat veel mensen denken, is dat niet Kalverstraat en Leidseplein in Amsterdam.

Een aantal hoofdstukken bevat een samenvatting en/of test-jezelf vragen. De antwoorden op de vragen staan achterin het boek vermeld. Een index is aanwezig. Het boek bevat veel aanknopingspunten om in een lessituatie het gebruik van kansrekening in het dagelijkse leven toe te lichten. Het staat inderdaad niet doorspekt met formules, zodat het de aandacht van formulehaters misschien langer vasthoudt. De leesbaarheid is prima. Niet alle tv-spelletjes waar de auteur over schrijft zijn in Nederland bekend, maar dat hoeft geen probleem te zijn. Veel voorbeelden die in het boek besproken worden zijn niet echt vernieuwend. De hoeveelheid bij elkaar, samen met een aantal minder vaak in dit soort boeken genoemde onderwerpen, zorgt ervoor dat het geheel aardig is om door te kijken.

*ir. M.M. Jansen
studierichting Bedrijfskader
Fontys Hogescholen*

Vapnik, V.N.

The nature of statistical learning theory

2000, Statistics for engineering and information science,
Springer-Verlag, Berlin,
300 blz., DM 139.00, ISBN 0-387-98780-0.

Via de VVS
25% korting!

Zoals de titel doet vermoeden, gaat dit werk over leertheorie. Wat moet men weten van een onbekende afhankelijkheid, om die te kunnen schatten met behulp van observaties? De laatste jaren zijn nieuwe methoden ontwikkeld om onbekende functies met een hoge dimensionaliteit goed te kunnen schatten met een relatief klein aantal observaties. Deze methoden hebben vorm gekregen in Support Vector Machines (SVM). SVM Margin control methods is nu een van de meest belangrijke richtingen in leertheorie en de toepassingen daarvan (voorwoord tweede druk).

Een aantal behandelde onderwerpen Na een korte inleiding waarin de geschiedenis van de leertheorie beschreven staat, definieert de schrijver het leerprobleem. Er wordt aandacht besteed aan noodzakelijk en voldoende voorwaarden voor convergentie. In het hoofdstuk over patroonherkenning komt de SVM voor het eerst ter sprake. De SVM-methode wordt gegeneraliseerd voor het schatten van reële functies. Directe leermethoden gebaseerd op het oplossen van meerdimensionele integraalvergelijkingen maken gebruik van SVM. Het empirical risk minimization principle en zijn toepassingen worden die richting uitgebreid.

Doel De schrijver beoogt de natuur van statistische leertheorie te beschrijven. Hij wil dat doen door dingen zo simpel mogelijk te beschrijven, zonder conceptuele versimpelingen. De details van de theorie en bewijzen zijn niet opgenomen. Hij verwijst daarvoor naar zijn andere boeken (*Estimation of dependencies based on empirical data* - Springer 1982 & *Statistical Learning theory* - Wiley 1988).

Doelgroep De schrijver geeft aan dat dit boek bedoeld is voor een brede doelgroep: studenten, ingenieurs en wetenschappers van verschillende achtergronden (statistiek, wiskunde, natuurkunde, informatici). Hij zegt dat het boek te begrijpen valt zonder kennis van speciale gebieden van de wiskunde. Hij zegt zich te realiseren dat het niet makkelijk lezen is.

Geschiktheid voor de doelgroep - Leesbaarheid Het lijkt mij dat dit boek geschikt is voor een smallere doelgroep dan Vapnik beschrijft. Het is redelijk taai lezen, gelukkig maken de "Informal reasoning and comments" aan het eind van elk hoofdstuk wel wat goed.

Sterke / Zwakke punten Het is een aardig idee om de bewijzen en details weg te laten om een beter leesbaar verhaal over te houden. Helaas vond ik het niet vlot leesbaar. De "Informal reasoning and comments" zijn beter te lezen. Het laatste hoofdstuk is zeer de moeite waard: "What is important in Learning theory?". En dan vooral de laatste paragraaf: "What is the most important?".

ir. M.M. Jansen
studierichting Bedrijfskader
Fontys Hogescholen

Wilkinson L.

The grammar of graphics

1999, Statistics and computing,

Springer-Verlag, Berlin,

xvii+408 blz, DM 139.00, ISBN 0-387-98774-6.

Via de VVS
 25% korting!

Eén plaatje zegt meer dan 1000 woorden. Wil dat cliché in de statistiek gelden, dan moet een plaatje goed zijn voor duizend getallen. En dat klopt: in de afgelopen jaren zijn goede grafieken sterke verkoopargumenten voor statistische software geworden. Dankzij Windows is er ook uniformiteit in het gebruik van printers gekomen en komt wat je op het scherm ziet meestal aardig overeen met wat op papier verschijnt. Ook is de bediening van de software meer en meer een grafische aangelegenheid geworden.

Voor veel grafieken geldt dat ze door de leverancier van de statistische software bedacht en geïmplementeerd zijn. De gebruiker krijgt nog enige vrijheid in het kiezen van kleur en type van lijnen en symbolen, maar daar is het wel mee gezegd. Wie een nieuw type grafiek bedenkt, of een wezenlijke variant op een bestaand type, zal vaak op

een zeer elementair niveau van grafische bewerkingen terug moeten vallen. Wilkinson claimt hier theoretisch en praktisch een oplossing voor te bieden. Theoretisch, door een grammatica van (statistische) grafieken te ontwerpen en te beschrijven; praktisch door (experimentele) software aan te bieden.

De grammatica is erg interessant: een grafiek wordt gezien als bestaand uit de volgende zeven onderdelen:

1. DATA: bewerkingen die gegevens uit een dataset halen en geschikt maken voor presentatie.
2. TRANS: transformaties, zoals wiskundige functies, hercodering en sorteren.
3. FRAME: een set variabelen (meestal 2 of 3) die de dimensies bepalen.
4. SCALE: schaaltransformaties, zoals logaritme of de vierkantswortel.
5. COORD: een coördinatenstelsel, bijv. rechthoekig of polair.
6. GRAPH: punten, lijnen, vlakken en hun eigenschappen, zoals kleur, raster en afmetingen.
7. GUIDE: assen en legenda.

Deze onderdelen worden gecombineerd met een instructies die sterk lijken op wat men in moderne object-georiënteerde programmeertalen (Delphi, Java, Visual Basic) ziet: eigenschappen van objecten worden gewijzigd, dan wel functies van objecten aangeroepen. Het boek bevat een enorm aantal voorbeelden van grafieken, de meeste in kleur en in een aantrekkelijke rustige stijl. Hun boodschap komt duidelijk over.

Of de grammatica correct en ondubbelzinnig is valt niet te controleren. Het boek staat vol met vele voorbeelden, maar levert geen tutorial: wie wil experimenteren is aangewezen op proberen van variaties. Daarbij is een complicatie dat de notatie een mengsel is van de zeven bovengenoemde onderdelen, in de vorm van labels en de object-georiënteerde aanroep van functies, waarvan de argumenten wisselen in aantal en positie. In de inleiding wordt gezegd dat een experimenteel systeem, GPL, beschikbaar zou zijn. De verwijzing naar de website van SPSS en SYSTAT - Wilkinson is en was architect en programmeur van laatstgenoemd pakket - loopt dood. Ook mijn speurtocht via zoekmachines op Internet leverde niets op. Er valt dus niet te experimenteren met Wilkinson's grammatica.

Wilkinson besteedt kort aandacht aan XML (Extensible Mark-up Language). Dat is een interessant verband met mogelijk vertrekkende mogelijkheden. XML de mogelijkheid de structuur van toegestane documenten te beschrijven (in een DTD, een Document Type Definition), in feite ook een soort grammatica. De presentatie kan met XSL (Extensible Style Language) gespecificeerd worden. Er bestaan al standaards voor DTDs van wiskundige formules en vectorgrafieken. Het zou fraai zijn als Wilkinson's grammatica in de vorm van een XML-standaard gegoten kon worden

- en door veel softwarebouwers overgenomen werd. Een DTD beschrijft dan de structuur van een grafiek, een XML-document bevat een daadwerkelijke grafiek en een XSL-document zou kleuren, lijndikten en zo specificeren. Achteraf aanpassen van de stijlaspecten van een grafiek (of een reeks grafieken) wordt zo eenvoudig. Hier ligt een zeer interessant terrein braak voor XML-enthousiasten.

Wilkinson stelt zich niet bescheiden op: hij heeft het over de grammatica van grafieken en spreekt die claim ook expliciet uit. Hij heeft over veel zaken een mening en overspeelt wel eens zijn hand. Zo noemt hij SQL consequent Sequential Query Language - een fraaie fout, ongetwijfeld veroorzaakt door de benaming Sequel die in bepaalde database-kringen als het watermerk van deskundigheid geldt. Wilkinson's uitstapjes naar de psychologie van de waarneming overtuigen allerm minst; ook doet zijn gewoonte hoofdstukken met een uitleg van de Griekse of Latijnse wortels van grafische termen te beginnen pedant aan.

Wat moet nu de conclusie zijn? Als aanzet tot denken over de structuur van grafieken is het boek zonder meer geslaagd. De talloze voorbeelden zijn zeer inspirerend. Sommige hoofdstukken, zoals dat over coördinatenstelsels, zijn bepaald sterk. Als gereedschap blijft de grammatica helaas vaporware: de aangekondigde software is onvindbaar en de beschrijvingen zijn te fragmentarisch om de praktische bruikbaarheid te kunnen beoordelen.

*P. Eilers
Sectie Medische Statistiek
Leids Universitair Medisch Centrum*

Ross, S.M.

Introduction to the mathematics of finance: options and other topics

1999, Cambridge University Press, Cambridge,
xv+184 blz, £ 21.95, ISBN 0-521-77043-2.

The target group of this book is defined by professional traders as well as undergraduates studying the basics of finance.

The first recommendation is to read this book just in one day or two consecutive days. Don't worry about what's happening at the stock- or option market!

The second recommendation is in order to get some experience to actually buy some stocks and options at the financial market and to sell them when the conditions of the Black-Scholes formula are met

Because this is the essence of the book: in chapter 7 the Black-Scholes option pricing formula is derived. Without the knowledge of the contents of the first six chapters it is impossible to understand this derivation.

The author of this book succeeded very well in explaining what should be understood of the conceptions probability, events, normal random variables, geometric Brownian motion (now I understand what's meant by this motion too!) (discretely as well as continuously varying) interest rates, and present value analysis. These topics one can read in chapters 1 to 4.

Chapters 5 and 6 pay attention to arbitrage, which is a sure-win betting scheme. Arbitrage can be used to determine the cost of an option. With a lot of examples the reader becomes familiar to the subject. The arbitrage theorem is presented, which states that exactly one of the following is true: either (a) there is a probability vector $\mathbf{p} = (p_1, p_2, \dots, p_m)$ for which

$$\sum_{j=1}^m p_j * r_i(j) = 0 \text{ for all } i = 1, \dots, n$$

$[r_i(\cdot)$ is the return function for a unit bet on wager i], or else (b) there is a betting scheme $\mathbf{x} = (x_1, \dots, x_n)$ for which

$$\sum_{i=1}^n x_i * r_i(j) > 0 \text{ for all } j = 1, \dots, m.$$

$[x_i$ is a bet on wager i]. This theorem is proved using the Duality Theorem of Linear Programming.

In chapter 7 then - already mentioned - the derivation of the Black-Scholes option pricing formula via the use of the multiperiod binomial model. The Black-Scholes formula is a formula of the unique cost of a call option that does not result in an arbitrage when the underlying security's price follows a geometric Brownian motion with volatility parameter sigma.

In the fourth paragraph of chapter 7 the author is dealing with American put options and gives an approximation of the risk-neutral price of an American put option. An American style put option allows the owner to put the stock up for sale at any time up to the expiration time of the option.

Chapter 8 handles the portfolio selection problem, the capital assets pricing model and mean variance analysis of risk-neutral-priced call options. The paragraph that handles mean variance analysis is, I think, too short.

Chapter 9 goes into exotic options like barrier options, Asian options and lookback options. A good method to price exotic options is by simulation. Several simulation estimators are discussed.

Chapter 10 is more practical: based on prices of crude oil data the statement is that these prices are not consistent with the assumption that crude prices follow a geometric Brownian motion. Another model is introduced.

Chapter 11 at last discusses autoregressive models and mean reversion.

And then suddenly the book appears to be out. I think an overall final comment could make the book still more readable.

Strong points of the book are the examples, the exercises, the good introduction of each paragraph, and the fact that each chapter has its own list of literature. Weak points are that the solutions of the exercises are not in the book as well, and that the ending of the book is too sudden.

Readers have to know the products of the stock-and optionmarket, and real experience is a prae. Although the author writes that it is not necessary to have knowledge of probability, I think it's advantageous to have it.

drs A.B.G.M. de Kroon

*Ordina Finance Business Intelligence Solutions BV
Amsterdam*

Beltrami, E.

What is random? Chance and order in mathematics and life

1999, Springer-Verlag, Berlin,

xx+201 blz., DM 39.90, ISBN 0-387-98737-1.

Via de VVS
25% korting!

Toeval is een onderwerp met genoeg intrigerende kanten om ook populariserende boeken over te schrijven, zoals moge blijken uit dit boek, waarin wiskundige Edward Beltrami aannemelijk probeert te maken dat orde en toeval twee kanten van dezelfde medaille zijn. Beltrami wil een breed publiek bereiken door allereerst te trachten een verhaal te vertellen, en dan voor wiskundige verdieping of precisering te verwijzen naar notenapparaat en appendices. Dit op zichzelf loffelijk streven is naar mijn smaak niet erg gelukt. Deskundigen op de diverse gebieden die hij bestrijkt (en dat zijn er nogal wat: wiskunde, communicatiewetenschappen, informatica, natuurkunde, filosofie, biologie, psychologie) zullen weinig nieuws zien en zich afvragen waar deze warrige gedachtegang toe moet leiden, terwijl de geïnteresseerde leek, alnaargelang zijn of haar karakter, ofwel nodeloos geïmponeerd raakt door de geëtaleerde technische kennis ofwel gefrustreerd achterblijft omdat de rondleiding wel nogal wat liet zien, maar nauwelijks verheldering heeft gegeven.

Het boek bevat de volgende hoofdstukken:

1. **The Taming of Chance** Een kleine geschiedenis van de waarschijnlijkheidsrekening die vooral bedoeld is om de lezer bekend te maken met de centrale limietstelling en de normale verdeling. De titel van het hoofdstuk is overigens zonder verwijzing geleend van Ian Hacking, wiens boek uit 1990 heel wat verrassender zaken aansneet. Adolphe Quetelet wordt hier ten onrechte afgeschilderd als een Franse astronoom wiens gemiddelde mens de grondslag is van de sociale wetenschappen.
2. **Uncertainty and Information** Hierin wordt de informatica geïntroduceerd als een additionele manier om toeval te begrijpen. De begrippen entropie, redundantie, en codering blijven het hele boek door een rol spelen. De preciezen onder ons zullen zich ergeren aan slordigheden zoals (op p. 39) dat er 1024 mogelijke woorden gevormd kunnen worden van lengte 10 bij een alfabet van 8 "even waarschijnlijk" symbolen.
3. **Janus-faced Randomness** Een poging om aan de hand van het genereren van random binaire reeksen aan te tonen dat "what appears as predictable is actually

randomness in disguise". Hier komt ook de Tweede Wet van de Thermodynamica aan de orde, waarin entropie in conflict zou komen met de reversibiliteit van de tijd. Daarover laat ik graag Bruno de Finetti aan het woord: "... it is often asserted - especially by philosophers - that the calculus of probability proves that the 'ergodic death' of the universe is inevitable. On the contrary, the calculus of probability (the logic of uncertainty) is completely neutral with respect to facts and behavior relating to natural phenomena (...) Ergodic death is very probable if one accepts (...) that model of physical phenomena which regards things as deriving from the destruction of an initial state of order (as in the mixing of gases; kinetic theory). But the calculus of probability in no way precludes phenomena in which a new order is created" (De Finetti, 1975, p. 150).

4. Algorithms, Information and Chance Hierin wordt algoritmische "randomness" en complexiteit aan elkaar gekoppeld in een erg gereduceerd, algebraïsch toevalsconcept, dat weinig statistici zal aanspreken ook al vallen er beroemde namen zoals Kolmogorov, Goedel, en Turing. Dit alles leidt tot de bewering dat elke poging om te besluiten of een lange binaire reeks random is of niet tot mislukken gedoemd is. Helemaal ten diepste beoordelen kan ik het niet, maar het lijkt me een rijkelijk overtrokken conclusie.

5. The Edge of Randomness Hierin wordt het thema van orde en chaos nog eens uit de doeken gedaan aan de hand van het bekende verhaal over fractals, en leren we naast 'white noise' ook 'pink noise', 'red noise' en 'black noise' kennen.

Al is het boek aan de oppervlakte vrij duidelijk geschreven, vanwege het gebrek aan conceptuele helderheid en vanwege de onbegrijpelijke pointe kan ik niemand aanraden het te kopen en van voren naar achteren te lezen. De literatuurlijst bevat wel een aantal interessante verwijzingen. De reden dat dit boek zo geforceerd aandoet is volgens mij dat toeval gedefinieerd wordt door niet alleen vormen van onafhankelijkheid aan te nemen, waar weinig statistici bezwaar tegen zullen maken, maar ook homogeniteit van kansen in dynamische processen. Wat zulke processen met gelijke kansen ook waard mogen zijn, ze komen in de natuur en het dagelijks leven uiteraard nauwelijks voor. Een veel sterker verhaal over de conceptualisering van toeval in de natuurkunde is te vinden in Von Plato (1994).

Verwijzingen:

- [1] B. De Finetti (1975). *Theory of Probability: A Critical Introductory Treatment* (Vol. 2). New York: Wiley.
- [2] I. Hacking (1990). *The Taming of Chance*. Cambridge: Cambridge University Press.
- [3] J. Von Plato (1994). *Creating Modern Probability: Its Mathematics, Physics and Philosophy in Historical Perspective*. Cambridge: Cambridge University Press.

W.J. Heiser

*Departement Psychologie, sectie M&T
Rijksuniversiteit Leiden*

Rao, C.R., Toutenburg, H.

Linear models. Least squares and alternatives

1999, Springer series in statistics,

Springer-Verlag, Berlin,

xvi+427 blz, DM 136.00, ISBN 0-387-98848-3.

Via de VVS
25% korting!

Dit boek gaat voor een groot deel over verschillende schatters voor de parameters van lineaire modellen. Deze schatters worden afgeleid uit verliesfuncties. Verdelingstheorie speelt een bescheiden rol. Het boek is onder meer bedoeld voor gevorderde studenten in de statistiek en verwante gebieden. Het eerste hoofdstuk geeft een beschrijving van de inhoud van het boek, in de volgende acht hoofdstukken staan regressie modellen voor continue variabelen centraal, en het laatste is gewijd aan modellen voor categorische variabelen. Het boek bevat een paar tabellen voor verdelingen en appendices over matrix algebra en software.

Hoofdstuk 2 geeft, na een kort inleidend voorbeeld van een economische vergelijking, in 13 bladzijden een overzicht van een aantal belangrijke klassen van lineaire regressie modellen, gedefinieerd door de structuur van de covariantiematrix van de fouten of storingen: achtereenvolgens stelsels simultane vergelijkingen, wat genoemd worden het multivariate en klassieke multivariate regressie model, en het gegeneraliseerde en standaard lineaire regressie model (met respectievelijk een algemene en scalaire covariantiematrix van de storingen).

De hoofdstukken 3 tot 5 behandelen parameter schatters in lineaire regressie modellen voor respectievelijk het standaard lineaire regressie model, het gegeneraliseerde model, en voor modellen met (stochastische) restricties voor de parameters. Het hoofdthema in deze hoofdstukken is het afleiden van schatters onder minimalisatie van verschillende verliesfuncties en, ook maar niet uitsluitend, randvoorwaarden als zuiverheid. Dit wordt gedaan voor schatters die lineaire functies zijn van de waarnemingen. Voor het vergelijken van schatters wordt de MDE (mean dispersion error) gebruikt of een lineaire functie daarvan. Voor een schatter b van een parametervector B is de MDE gedefinieerd als $\mathcal{E}(b - B)(b - B)'$, waarin b niet zuiver hoeft te zijn.

Nadat in hoofdstuk 3 aldus schatters zijn afgeleid en vergeleken, worden meest aannemelijke schatters besproken met daarop aansluitend toetsing van en betrouwbaarheidsintervallen voor lineaire functies van de parameters. Vervolgens komen in dit hoofdstuk een reeks onderwerpen aan de orde. Een aantal wordt vrij kort behandeld (principale componenten regressie, krimp schatters en partial least squares als methoden voor het schatten onder multicollineariteit, projection pursuit regressie, censuring, schatten van nonparametrische en logistische regressies met neurale netwerken). Meer aandacht wordt besteed aan ridge regressie in de context van multicollineariteit, aan het minimax principe voor het schatten onder ongelijkheidsvoorwaarden en aan lineaire minimax schatters. De verliezen van de ridge schatters en de lineaire minimax schatters worden vergeleken met die van de OLS schatters en de voorwaarden waaronder de eerste twee superieur zijn worden afgeleid. Bovendien worden de lineaire minimax schatters geïnterpreteerd als ridge schatters.

In hoofdstuk 4 worden uit verschillende verliesfuncties lineaire schatters afgeleid voor het regressie model met een algemene covariantiestructuur. Dit leidt tot schatters die niet steeds zuiver hoeven te zijn, maar natuurlijk komt ook de beste lineaire zuivere schatter aan bod. Als voorbeelden van deze modellen worden heteroscedasticiteit en autoregressie van de storingen gegeven. In relatief kort bestek worden gemengde modellen behandeld. De presentatie is gebaseerd op vrij recent werk van Rao, is zeer algemeen en fraai, maar daardoor zal de naive lezer niet gauw de indruk krijgen dat het schatten van variantiecomponenten of stochastische effecten geen eenvoudig probleem is, dat nog steeds niet in het algemeen is opgelost. Het hoofdstuk besteedt ook wat ruimere aandacht aan empirische Bayes schatters voor een hiërarchisch model en natuurlijk wordt het verlies vergeleken met dat van de OLS schatter. Het hoofdstuk besluit met een korte behandeling van GLS schatters onder verschillende condities, waaronder die van een singuliere covariantiematrix, gebaseerd op al wat ouder en zeer bekend werk van Rao.

Het gebruik van a priori informatie is het onderwerp van hoofdstuk 5. Allereerst komt natuurlijk het schatten onder exacte lineaire restricties aan de orde. Daarbij wordt de kwaliteit van schatters onder twee verschillende restricties vergeleken. Ook worden schatters vergeleken waarvoor de restricties niet exact juist zijn (de onjuistheid mag al of niet stochastisch zijn, of een combinatie hiervan) en die daardoor onzuiver zijn. Dit systematische overzicht steunt op een paar algemene stellingen over voorwaarden waaronder verschillen tussen matrices niet negatief definitief zijn. De resultaten worden toegepast op lineaire, mogelijk onzuivere schatters. Vervolgens worden stochastische restricties behandeld en schatters die wat genoemd wordt 'weakly (R, r) -unbiased' zijn (een voor mij nieuw begrip).

Hoofdstuk 6 behandelt voorspellen in regressie. Het Kalman filter wordt hier afgeleid als een speciaal geval van een lineaire voorspeller met minimale gemiddelde dispersie onder verschillende vormen van a priori informatie. Hoofdstuk 7 gaat onder de naam sensitiviteitsanalyse over de relatieve invloed van verschillende waarnemingen op de schatters, wat toepassing vindt in hoofdstuk 8, waarin het schatten bij ontbrekende waarnemingen behandeld wordt. Hoofdstuk 9 geeft in 17 bladzijden een overzicht van robuuste regressie (least absolute deviations schatters, M-schatters en hun asymptotische verdeling). In dit hoofdstuk worden geen bewijzen van de gepresenteerde resultaten gegeven. Hoofdstuk 10 tenslotte behandelt modellen voor categorische variabelen in de context van de gegeneraliseerde lineaire modellen van Nelder en Wedderburn en hun (quasi) meest aannemelijke schatters. Als voorbeelden dienen logistische en loglineaire modellen. In het laatste deel van dit hoofdstuk komen modellen ter sprake voor gecorreleerde categorische waarnemingen. Daarbij wordt gekozen voor de GEE benadering van Liang en Zeger. Ook hier worden eigenschappen van de schatters wel gegeven maar niet bewezen.

Wie een boek over lineaire modellen besproken ziet waarvan C.R. Rao een van de auteurs is, wil misschien weten hoe dit boek zich verhoudt tot Rao's vermaarde *Linear Statistical Inference and Its Applications*. Het hier besproken boek is minder (abstract) wiskundig. Behalve in de laatste twee hoofdstukken bestaat het gebruikte gereedschap

voornamelijk uit verwachtingen en matrix algebra. Verdelingstheorie speelt hier geen belangrijke rol. De appendix bevat een aantal verdelingstellingen, maar die worden niet bewezen, tenzij als gevolg van een eerdere stelling die niet bewezen wordt.

Toch lijkt het boek mij vooral aan te bevelen voor studenten en lezers met (ook) een theoretische belangstelling. Zij vinden in de eerste zes hoofdstukken en de appendix over matrix algebra een grote hoeveelheid resultaten bijeengebracht, die soms zeer algemeen en elegant worden afgeleid. Dat wil niet zeggen dat de volgorde waarin de zaken gepresenteerd worden mij altijd overtuigt. Het is soms lastig om door de bomen het bos te blijven zien, waarschijnlijk ook omdat sommige onderscheidingen nogal subtiel zijn.

Ik denk niet dat lezers die weinig ervaring hebben met lineaire modellen door dit boek een gevoel krijgen voor de verscheidenheid aan situaties waarin varianten van deze modellen gebruikt (kunnen) worden. Daarvoor krijgen niet alleen toepassingen te weinig aandacht, maar wordt volgens mij ook te weinig aandacht gegeven aan de context en overwegingen die leiden tot het gebruik van bepaalde modellen en schatters. Variantie-analyse wordt bijvoorbeeld niet genoemd als een voorbeeld van modellen waarin de matrix van regressoren singulier is. De lezer die niet beter weet zal ook niet gauw een indruk krijgen van de verscheidenheid aan covariantiestructuren van de storingen die zich kunnen voordoen. Als leerboek zou het boek misschien beter geslaagd zijn als aan zulke zaken meer aandacht was gegeven, desnoods ten koste van andere onderwerpen die toch maar heel summier behandeld worden. Dat merk je vooral als je weinig van een onderwerp weet. In hoofdstuk 9 worden de voorwaarden waaronder robuuste schatters hun eigenschappen hebben niet of nauwelijks toegelicht. Je krijgt dan de indruk dat deze voorwaarden vooral gekozen zijn om tot het gewenste resultaat te kunnen komen. Daardoor lijken sommige formules al gauw bijna even deerniswekkend als King Kong in zijn laatste ogenblikken.

A.L. Beem

*Afdeling Biologische Psychologie
Vrije Universiteit Amsterdam*

Efromovich, S.

Nonparametric curve estimation

1999, Springer series in statistics,

Springer-Verlag, Berlin,

425 blz, DM 148.00, ISBN 0-387-98740-1.

Via de VVS
25% korting!

De verschillende statistische probleemstellingen, kansdichtheidsschatten, nonparametrische regressie, tijdreeks-analyse en spectraal-analyse, en filtering van stochastische processen (met één tijdsparemeter) worden geanalyseerd met één overkoepelende methode, namelijk met behulp van orthogonale funktiesystemen, zoals Fourierreeksen en wavelets. De schatting wordt hoofdzakelijk aangedreven door de data. Om tot redelijke keuzes voor de parameters in zulke methodes te komen, heeft de auteur een

8-tal *corner functions* gepostuleerd, referentiefuncties op het eenheidsinterval, in termen van dichtheden variërend van de uniforme dichtheid via de afgekapte standaard-normale dichtheid, twee bimodale gladde dichtheden, een afgekapte zeer geconcentreerde normale dichtheid, een niet-gladde dichtheid en een monotone dichtheid tot een discontinue dichtheid. De keuzes van parameters in de schattingsmethodes worden gerechtvaardigd op grond van simulaties, waarbij de hierboven vermelde referentiefuncties als te schatten object dienst doen.

Er is een *online* beschikbaar S-PLUS pakket, waarmee de grafieken in het boek gereproduceerd kunnen worden, en de afhankelijkheid van de parameters nader bekeken kan worden. Er zijn de praktijkseminaria (begeleiding van uit te voeren onderzoek van een gegeven dataset), casus-studies (analyses van meer specifieke modellen) en keuzesonderwerpen door het hele boek heen, en opgaven aan het eind van ieder hoofdstuk.

Het grootste gedeelte van het boek is nogal praktisch gericht. De nadruk ligt op de analyse van kleine datahoeveelheden. Ik veronderstel, dat de toepasser in zijn toepassingsgebied een versie van de referentiefuncties maakt en de voorgestelde schattingsmethodes volgens de overwegingen in het boek zo aanpast dat ze naar tevredenheid werken bij zijn referentiefuncties. Er wordt echter ook gewerkt aan inzicht. Vanuit theoretisch standpunt het meest bevredigend uitgewerkt is een hoofdstuk waar de asymptotiek bekeken wordt van datagestuurde schattingsmethodes, een specialiteit van de auteur. In dat hoofdstuk wordt een analogie aangegeven tussen de verschillende probleemstellingen, die verder gaan dan de overeenkomst in methode. In het bijzonder wordt er een zogenaamd *principle of equivalence* aangestipt, dat zegt dat resultaten betreffende risico's met een begrensde verliesfunctie in het filteringmodel onder milde aannames ook geldig zijn in de andere in dit boek besproken probleemstellingen. Het laatste hoofdstuk is gewijd aan schattingsmethodes die niet expliciet berusten op orthogonale functiesystemen.

Het is een boek dat prettig, stimulerend en uitdagend is om uit te lezen en te werken. Men vindt er allerlei verfrissende voorbeelden en tegenvoorbeelden.

*dr H.W.M. Hendriks
Subfaculteit Wiskunde
Katholieke Universiteit Nijmegen*

Serfozo, R.

Introduction to stochastic networks

1999, Applications of mathematics,
Springer-Verlag, Berlin,
330 blz., DM 139.00, ISBN 0-387-98773-8.

Via de VVS
25% korting!

The title of this book is a misleading one: A novice in queueing or stochastic networks who tries to learn about this subject using this book will have a hard time. Most results are presented in their most general form, and many traditional results (such as the stationary distribution of Jackson networks, Poisson output of birth-death processes, etc.)

are only treated as examples of more general results. This is an interesting approach for readers already familiar with queueing networks, but the level is too high for an introduction.

This focus on depth is also seen in the subjects treated. They consist mainly of those related to Markov processes and point processes. Many other subjects (such as general service times, controlled networks, mean value analysis, etc.) are not or hardly discussed. This reflects of course the research interests of the author. Also for this reason this text is not suitable for an introduction to the subject.

The above does not mean that this book has no merit. It is a very well written overview of the state-of-the-art. Advanced subjects are treated in depth and with clarity.

The first chapters deal with "standard" queueing theory: (generalizations of) Jackson and BCMP networks, reversibility, sojourn times, Palm probabilities, Little laws. Later chapters are concerned with other types of stochastic networks, such as those with string transitions, a special class of quasi-reversible nets, and spatial queueing systems.

*G.M. Koole
Faculteit der Exacte Wetenschappen
Vrije Universiteit Amsterdam*

Leonard, T., Hsu, J.S.J.

Bayesian methods: an analysis for statisticians and interdisciplinary researchers

1999, Cambridge series in statistical and probabilistic mathematics,

Cambridge University Press, Cambridge,

xiv+333 blz, £ 40.00, ISBN 0-521-59417-0.

Het materiaal in dit boek is door de auteurs gebruikt bij het onderwijs aan doctoraal studenten in wiskunde, economie en bedrijfskunde. Behalve voor deze doelgroep is het boek volgens de auteurs ook bedoeld voor onderzoekers uit diverse vakgebieden, die Bayesiaanse methoden in hun vakgebied willen gebruiken, en voor studenten (en onderzoekers) op het gebied van statistische methoden. De laatste groep zou met behulp van de hoofdstukken 5 en 6 de huidige grenzen kunnen bereiken van het onderzoek in Bayesiaanse statistiek.

Het boek begint met een gedegen introductie van klassieke Fisherianse aanpak. Aan bod komen, onder andere, likelihood principe, sufficiency, profile likelihood en EM-algoritme. Hoofdstuk 2 behandelt de discrete versie van de stelling van Bayes, die onder andere wordt toegepast bij het probleem van model-selectie. Hoofdstuk 3 is gewijd aan continue variabelen met 1 onbekende parameter. Hierbij krijgen ook de frequentistische eigenschappen van de Bayesiaanse procedures de aandacht. Hoofdstuk 4, over utilities, staat tamelijk los van de rest van het boek. Hoofdstuk 5 behandelt problemen met meerdere onbekende parameters. De Bayesiaanse aanpak vereist het vinden van de marginale a posteriori verdeling. Dit lastige probleem kan op verschillende

manieren aangepakt worden. Het boek legt de nadruk op de Laplaciaanse benadering. Behandeld worden onder andere de normale lineaire regressie, de logistische regressie en de analyse van tijdreeksen. In het laatste hoofdstuk, "Prior Structures, Posterior Smoothing, and Bayes-Stein Estimation", komen diverse topics aan bod. Onder andere hiërarchische modellen en computationele aspecten, waaronder - heel summier - Markov Chain Monte Carlo (MCMC) benadering.

De presentatie van het materiaal zit heel goed in elkaar. Theoretische paragrafen geven, na de uitleg van de theorie, een of meer uitgewerkte voorbeelden (Worked Examples voorzien van Model Answers). Daarna volgen opgaven voor zelfstudie. Door middel van deze opgaven wordt de gepresenteerde theorie aangevuld, verdiept en toegepast in een meer specifieke context. Vaak is het onderwerp van een opgave oorspronkelijk in een afzonderlijke artikel beschreven; soms vereist het maken van de opgave het raadplegen van zo'n publicatie. Op de via opgaven geïntroduceerde theorie wordt in de latere hoofdstukken voortgebouwd, de lezer kan de meeste opgaven derhalve niet overslaan.

Het boek bevat 49 uitgewerkte voorbeelden en 148 opgaven zonder uitwerkingen. Het spectrum van besproken applicaties is zeer breed. Ik ben enkele drukfouten tegengekomen, maar over het geheel is het boek bijzonder goed verzorgd.

Ondanks de vele toepassingen die het boek bespreekt, denk ik dat het voor de "interdisciplinary researchers" niet gemakkelijk toegankelijk zal zijn. Bovendien hebben zulke onderzoekers waarschijnlijk vooral behoefte aan praktische toepassingen met behulp van de MCMC technieken, en die komen in het boek nauwelijks ter sprake. Het boek is zonder meer aan te raden aan mathematisch statistisch geschoolde lezers die actief (willen) zijn op het gebied van onderzoek of onderwijs in Bayesiaanse statistiek.

dr V. Fidler

*Vakgroep Medische Statistiek
Rijksuniversiteit Groningen*

Coxon, A.P.M.

Sorting data: Collection and analysis

1999, Sage, Thousand Oaks, California,
vi+98 blz, £ 8.99, ISBN 0-8039-7237-7.

Elke docent krijgt af en toe te maken met studenten die hun studiemateriaal, waaronder hele boeken, zelf gefotokopieërd hebben en – om de kosten nog meer te drukken – de pagina's met de literatuurvermeldingen niet hebben meegekopieërd. Nu is het één ding om een 'arme' student duidelijk te moeten maken dat dit een verkeerde, in ieder geval onprofessionele, vorm van zuinigheid is, maar om hetzelfde te moeten uitleggen aan een uitgever van wetenschappelijke boeken is ronduit verbijsterend. Vanwaar deze tirade?

De auteur A.P.M. Coxon, *senior lecturer* aan de *University of Essex*, onder andere bekend van een groot beroepenonderzoek (Coxon en Jones, 1978, 1979) en van een

boek over meerdimensionale schaaltechnieken (Coxon, 1982), heeft als deel 127 in de *Quantitative Applications in the Social Sciences* van Sage een boek geschreven dat zonder twijfel het predikaat 'most complete handbook around on sorting techniques' verdient dat de *Series Editor* (p. vi) aan dit boek toekent. Voor de tekst van dit ultieme boek over sorteermethoden heeft de auteur 97 bladzijden nodig gehad, en omdat hij kennelijk ook niet meer dan dat aantal bladzijden tot zijn beschikking kreeg, zijn de 65 voetnoten, de literatuurlijst met 265 bronnen, en de vier appendices niet in het boek opgenomen maar op een *internet*-site van Sage 'gepubliceerd'. Dat leidt ertoe dat de lezer steeds heen en weer moet pendelen tussen de basistekst en een eigen afdruk van een internetdocument van 25 bladzijden. Dat komt de handzaamheid niet echt ten goede.

Maar – het moet gezegd worden – ondanks deze knulligheid, is Coxons boek inderdaad het meest uitgebreide en gedetailleerde boek over de sorteermethode, met name over het verzamelen, beschrijven en analyseren van sorteergegevens. Sorteergegevens bestaan uit een of meer partitioneringen van een verzameling 'objecten'. Zulke partitioneringen kunnen onder andere verkregen worden als resultaat van de sorteermethode. In zijn simpelste vorm komt die hierop neer dat men een of meer respondenten een verzameling kaartjes voorlegt waarop de namen of afbeeldingen van de te beoordelen 'objecten' zijn weergegeven. De respondenten wordt nu gevraagd de kaartjes op stapeltjes te leggen, en wel zodanig dat kaartjes met objecten die 'bij elkaar horen' of 'op elkaar lijken' op hetzelfde stapeltje gelegd worden en kaartjes met objecten die 'niet bij elkaar horen' of 'niet op elkaar lijken' op verschillende stapeltjes terecht komen. Meestal mogen de respondenten net zoveel stapeltjes maken als zij zelf willen en net zoveel kaartjes op elk stapeltje leggen als zij vinden dat nodig is. Het resultaat is dat de objecten door elke respondent in elkaar uitsluitende categorieën van een nominale variabele worden ingedeeld. Zulke indelingen kunnen natuurlijk ook gevonden worden zonder dat men de objecten door respondenten laat sorteren. Men denke bijvoorbeeld aan de indeling van een verzameling boeken in de catalogi van verschillende bibliotheken.

Coxons boek bestaat uit vier hoofdstukken: 1. inleiding, 2. het verzamelen van sorteergegevens, 3. het beschrijven en vergelijken van sorteergegevens, en 4. het analyseren van sorteerddata. In het tweede hoofdstuk worden verschillende varianten van bovengenoemde basismethode besproken, en wordt aandacht besteed aan manieren waarop men objecten kan vinden die het onderzoeksdomein representeren. Daarnaast bespreekt Coxon verschillende eigenschappen van sorteercriteria (gelijkenis, bij elkaar horen) en geeft hij gedetailleerde aanwijzingen voor de instructie van de respondenten en de procedure van afname. Tot slot van dit hoofdstuk benadrukt hij de wenselijkheid van een pretest en suggereert hij een *preferred data format* voor de registratie van sorteergegevens.

In het derde hoofdstuk behandelt Coxon verschillende maten die de eigenschappen van afzonderlijke sorteerresultaten beschrijven en die het mogelijk maken de sorteringen van verschillende respondenten met elkaar te vergelijken. Een belangrijk kenmerk van een bepaalde partitionering is de 'hoogte', ook wel de 'klonterigheid' genoemd.

Hierbij gaat het om het aantal partities binnen de sortering en het aantal objecten binnen elke partitie: zijn er slechts een paar groepen met elk veel objecten of zijn er veel partities die stuk voor stuk weinig objecten bevatten? Om de overeenkomst tussen twee sorteringen vast te stellen zijn er verschillende indices ontwikkeld die berusten op het aantal paren van objecten die zowel bij de ene respondent als bij de andere respondent samen in een partitie terecht zijn gekomen. Coxon behandelt enkele daarvan in zijn hoofdstuk en geeft formules voor een aantal andere in een van zijn internet-appendices.

In Hoofdstuk 4, over het analyseren van sorteergegevens, gaat het om het gezamenlijk representeren van de sorteringen van verschillende respondenten. Hierbij kan de nadruk vallen op een representatie van de respondenten, op een weergave van de objecten of om een gezamenlijke afbeelding van respondenten en objecten. De methoden die Coxon bespreekt zijn meerdimensionale schaaltechnieken en clusteranalyse (en in samenhang daarmee het cruciale probleem hoe een adequate afstandsmaat te definiëren), correspondentie-analyse en drie vormen van multi-pele correspondentie-analyse. Dit laatste gebeurt wat onevenwichtig: van het programma MDSORT wordt in detail de verliesfunctie en de schattingsmethode in matrixnotatie weergegeven, terwijl de programma's HOMALS en MSA (multidimensionale scalogramanalyse) alleen op verbaal-abstrakte manier besproken worden, overigens zonder in te gaan op de verwantschap tussen deze drie methoden. Er wordt niet opgemerkt dat MDSORT identiek is aan HOMALS en dat MSA en HOMALS (en dus ook MDSORT) alle drie programma's zijn voor principale-componentenanalyse van nominale variabelen.

In elk hoofdstuk worden relevante (reken)voorbeelden behandeld, de ene keer wat gedetailleerder dan de andere keer. Dit laatste is overigens niet echt een bezwaar omdat de uitgebreide literatuurlijst (oepe!) voldoende verwijzingen bevat naar publicaties die de lezer die meer informatie wil kan raadplegen. Kortom, dit is een uiterst waardevol boek voor iedereen die wil kennismaken met de wereld van sorteergegevens. Het is de verdienste van Coxon dat die kennismaking bepaald niet aan de oppervlakte blijft. Voor degenen die nog dieper op het onderwerp in willen gaan, vormt Coxons boek bovendien een uitstekende wegwijzer in de uitgebreide literatuur. Het ligt dus voor de hand om dit boek te gebruiken bij onderwijs in de afstudeerfase aan studenten die vanuit wat voor achtergrond dan ook in onderzoek met de sorteermethode geïnteresseerd zijn.

Verwijzingen:

- [1] Coxon, A.P.M. (1982). *The user's guide to multidimensional scaling*. London: Heinemann.
- [2] Coxon, A.P.M., Jones, C.L. (1978). *The images of occupational prestige: A study in occupational cognition*. London: Macmillan.
- [3] Coxon, A.P.M., Jones, C.L. (1979). *Measurements and meanings: Techniques and methods of studying occupational cognition*. London: Macmillan.

dr W.A. van der Kloot
 Departement Psychologie
 Universiteit Leiden

Kallianpur, G., Karandikar, R.L.

Introduction to option pricing theory

1999, Birkhäuser, Basel,

280 blz, DM 128.00, ISBN 3-7643-4108-4.

Het boek gaat over stochastische processen in de wiskundige theorie van financiering, met name over optieprijzen. Uitdrukkelijk blijven futures en termijnstructuren buiten beschouwing.

Het boek bevat een introductiegedeelte over stochastische analyse (Itô's theorie, integratie over semimartingalen en stochastische differentiaalvergelijkingen) zodat voorkennis op dit gebied niet nodig is. Daarna volgen achtereenvolgens: optieprijzen in discrete tijd, continue tijd. Complete markten en vervolgens Europese, Amerikaanse en Russische opties (opties met een - eventueel - oneindige tijd tot uitoefening). Het boek is geschikt voor zelfstudie en voor gebruik bij een vak op doctoraal niveau.

Het boek kent de volgende hoofdstukken:

1. Stochastic Integration
2. Ito's Formula and its Applications
3. Representation of Square Integrable Martingales
4. Stochastic Differential Equations
5. Girsanov's Theorem
6. Option Pricing in Discrete Time
7. Introduction to Continuous Time Trading
8. Arbitrage and Equivalent Martingale Measures
9. Complete Markets
10. Black and Scholes Theory
11. Discrete Approximations
12. The American Options
13. Asset Pricing with Stochastic Volatility
14. The Russian Options

In het boek wordt verder geen aandacht besteed aan numerieke methoden, niet aan exotische opties en - zoals eerder gezegd - niet aan interestderivaten. Het boek bevat geen oefeningen. Wel heeft het een hoge formuledichtheid met veel aandacht voor bewijsvoering. Het boek is erg zorgvuldig geschreven (afgezien van een klein

typfoutje op pagina 106 heb ik verder geen fouten kunnen vinden). Het is niet zo praktijkgericht als bijvoorbeeld Willmot, Howison en Dewynne (1995). Het boek is eerder voor wiskundigen bedoeld dan voor direct gebruik in de praktijk. Het werk lijkt me wel interessant voor praktijkbeoefenaren die dieper willen duiken in de aannamen achter de formule van Black-Scholes en in de gevoeligheden van de formule voor deze onderliggende aannamen.

Verwijzingen:

- [1] Willmot, P., Howison, S. and Dewynne, J. (1995). The mathematics of financial derivatives, A student introduction. Cambridge, Cambridge University Press.

Th. Beekman

Le Gall, J.-F.

Spatial branching processes, random snakes and partial differential equations

1999, Lectures in mathematics,

Birkhäuser, Basel,

176 blz, DM 44.00, ISBN 3-7643-6126-3.

This book is based on a series of lectures given at ETH Zurich. The book gives an self-contained exposition of a part of probability theory that has been studied extensively in recent years. This text contains a presentation of superprocesses and their basic properties.

Consider a Galton-Watson process which describes the evolution of a population where at integer time points each particle in the population produces independently offspring according to some given law. A Galton-Watson process is an example of a branching process. A continuous-state branching process is a Markov process which can be obtained as a weak limit of a sequence of rescaled (in time and space) Galton-Watson processes. A continuous-state branching process can be characterized by a function ψ which is called the branching mechanism. The branching mechanism is a function of a special type determined by three parameters. This is basically the work of John Lamperti (1967). Assume now that the particles of the approximating sequence of Galton-Watson processes perform independently a motion in $\mathbf{R}^d$ (or more general in some Polish space E) according to the law of some given diffusion process ξ . The processes are considered as measure-valued where the state of the process is a random measure on $\mathbf{R}^d$ with support the set of positions of the particles. Under some conditions we get a weak limit. The limit is a spatial branching process which is characterized by the diffusion ξ and the branching mechanism ψ of the total mass. The limit process is also called a (ξ, ψ) -superprocess. Important is the connection between (ξ, ψ) -superprocesses and partial differential equations associated with the operator $Lu - \psi(u)$, where L is the generator of the diffusion ξ . If the diffusion is Brownian motion one gets super-Brownian motion which is important in models from statistical mechanics and infinite particle systems. Of course, an initial motivation for the study of superprocesses was the modelling of spatial populations, especially models with interactions (catalytic superprocesses).

The book is intended for postgraduate students and researchers in probability theory. As a prerequisite is assumed some knowledge of stochastic processes, especially branching processes (Galton-Watson process) and Brownian motion. Since the book is self-contained and very clearly written, with references to results that are not derived, it is to my opinion also accessible to readers with interest in developments in modern probability theory. In any case, one should read the first chapter which is a non-technical introduction to the results described in the book.

*J.A.M. van der Weide
Informatie Technologie en Systemen
Technische Universiteit Delft*

Padberg, M.

Linear optimization and extensions

1999, Algorithms and combinatorics,
Springer-Verlag, Berlin,
501 blz, DM 169.00, ISBN 3-540-65833-5.

Via de VVS
25% korting!

Dit boek is een tweede, gereviseerde en enigszins uitgebreide druk van de versie die in 1995 is verschenen. Het is op zich een verheugend feit dat een leerboek als dit een tweede druk bereikt. Het is namelijk een boek dat de student stevige kost voorziet. De behandeling van het hoofdonderwerp lineaire programmering is gedegen, dat wil zeggen dat de nodige bewijzen compleet vermeld worden.

Aan de orde komen de verschillende oplosprincipes voor lineaire programmering, te weten de simplexmethode, projectieve algoritmen en ellipsoid-methoden. Een hoofdstuk over analytische meetkunde behandelt de nodige voorkennis die niet elementair is.

De in de titel beloofde uitbreidingen betreffen de behandeling van combinatorische optimaliseringsproblemen zoals het Handelsreizigersprobleem; men zou dit hoofdstuk kunnen beschouwen als een illustratie van toepassingen van lineaire programmering.

De opgaven die het boek bevat zijn niet echt eenvoudig, zodat de aankondiging van een uitwerkingenboek zeer welkom is.

In een appendix staan tenslotte drie grote opdrachten die van real life omvang zijn en waarop de student zijn oploscapaciteiten kan laten zien.

Naar mijn mening is het boek uitstekend geschikt voor een cursus lineaire programmering voor een publiek dat niet bang is voor het gebruik van de nodige wiskunde.

De docent die het boek kiest moet zich er overigens wel van bewust zijn dat het jaar van publiceren 1999 niet betekent dat ook de laatste ontwikkelingen zijn meegenomen: de verwerkte literatuur is niet bij de tijd gebracht bij de revisie. Wellicht kan dat nog eens gebeuren bij de hopelijk nog te verschijnen derde editie, want die verdient dit boek zeker.

*dr A. Volgenant
Onderzoeksgroep Operationele Research en Management
Universiteit van Amsterdam*

Tayur, S., Ganeshan, R., Magazine, M. (eds)
Quantitative models for supply chain management

1998, International series in operations research and management science 17,
 Kluwer, Dordrecht,
 896 blz, fl 510.00, ISBN 0-7923-8344-3.

Since the early 90's, supply chain management has been a rapidly growing area of research, both in academia as well as in industry. Nowadays, there is an overwhelming body of literature that may be classified under this subject. As the title indicates, this book focuses on quantitative aspects of supply chain management. The aim of the editors was to provide, "in a unified manner, a systematic summary of the large variety of new issues being considered, the new set of models being developed, the new techniques for analysis, and the computational methods that have become available recently." This has resulted in a book with the impressive number of 26 chapters, many of them authored by well-known experts. All contributions were invited and had to fit into one of six categories preselected by the editors. The specific chapter topics were chosen by the authors themselves.

The categories in which the chapters are organized are "Basic Concepts and Technical Material" (Chapters 1-6), "Supply Contracts" (Chapters 7-10), "Value of Information" (Chapters 11-15), "Managing Product Variety" (Chapters 16- 19), "International Operations" (Chapters 20-22) and "Conceptual Issues and New Challenges" (Chapters 23-26). Some of the chapters present an overview of a specific research topic or they provide a framework for analysis, while others present primarily new results. In the latter case, the material is sometimes closely related to technical reports and journal articles that still had to appear in 1998 (when the book was published). In this way, the book presents quite a broad and up-to-date picture of quantitative supply chain management research.

One could ask, however, whether the editors have been successful in presenting this picture as clearly as possible. Besides their role in selecting the categories and authors, they seem to have confined themselves to writing a short introduction, which mainly consists of summarizing the chapters. In our opinion, the book could have benefited from more editing and a different organization. There are several annoying mistakes such as typos, missing or incorrect references, and some of the material is unnecessarily repeated. More importantly, the order of the chapters does not always seem to be the most logical one. For instance, within the categories one would expect the overview chapters to appear first, such that the reader gets a good idea of the relevant issues and main results before (s)he arrives at the chapters which present recent results on the topic. In most categories, a different order has been chosen.

A similar remark applies to the overall order of the chapters. The last chapter presents a very useful taxonomic review of supply chain management research, co-authored by one of the editors. It also discusses questions defining the field. "What is a supply chain?", "What does managing the supply chain mean?" and "Who manages the supply chain?". In our opinion it would have been much more logical to discuss

these questions in the introduction of the book (if only to point out that different people may give different answers). Moreover, given the fact that the editors were striving for a "systematic summary", provided in a "unified manner", it could have been useful to apply the proposed taxonomy to the chapters in this book.

Despite these critical remarks, we think that this book contains a lot of interesting material that is suitable as a graduate text and as a reference for researchers. Because of its price, however, some people may be reluctant to buy this book, especially because some of the journal articles related to some of the chapters have appeared in the mean time.

*D. Romero Morales
dr A.P.M. Wagelmans
Erasmus Universiteit Rotterdam*

Hooley, G., Hussey, M.

Quantitative methods in marketing

1999, International Thomson Business Press, London,
300 blz, £ 27.99, ISBN 1-86152-417-X.

Kwantitatieve modellen worden steeds populairder binnen marketing management. Ook het bedrijfsleven is zich bewust dat de marketing beslissingen moeten worden ondersteund met kwantitatieve analyses (naast de kwalitatieve analyses). Er is dus behoefte aan een overzicht van de belangrijkste en nieuwste technieken binnen de kwantitatieve marketing analyse.

In 1994 verscheen een Special Issue van Journal of Marketing Management (Vol 10, 1994) met als titel "Quantitative methods in marketing". Het doel van dit issue was om in een monograph vorm dit gebied in beeld te brengen. Dit boek is de 2e editie van het Special Issue uit 1994, met vernieuwde bijdragen uit de 1e editie en enkele nieuwe bijdragen. Omdat het Journal of Marketing Management sterk gericht is op de Europese markt, zijn opnieuw de meeste bijdragen afkomstig van Europese (meest Engelse) auteurs. De bijdrage van Scott Armstrong en Roderick Broderie met een overzicht van voorspeltechnieken in marketing is voor mij de belangrijkste verbetering ten opzichte van de eerste editie.

Dit boek verzamelt verschillende basistechnieken voor kwantitatief marketing management in deel 1, en illustreert deze vervolgens met een aantal casestudies in deel 2. De besproken onderwerpen zijn:

1. Quantitative Methods in Marketing: The Multivariate Jungle Revisited (Graham J. Hooley and Michael K. Hussey)
2. The Diffusion of Quantitative Methods into Marketing Management (Michael K. Hussey and Graham J. Hooley)
3. Cluster Analysis (John Saunders)

4. Logit Model Analysis for Multivariate Categorical Data (David Jobber)
5. Solving Marketing Problems with Conjoint Analysis (Marco Vriens)
6. Forecasting for Marketing (J. Scott Armstrong and Roderick J. Brodie)
7. The AHP: Structuring Marketing Information for Decision Support (Mark A.P. Davies)
8. The Structure of Consumers' Place Evaluations (Paul M.W. Hackett and Gordon R. Foxall)
9. An Introduction to Hierarchical Moderated Regression Analysis (Gordon E. Greenley)
10. Attitude Survey Data Reduction Using CHAID: An Example in Shopping Centre Market Research (Steve Baron and Dianne Philips)
11. Regression Type Techniques and Small Samples: A Guide to Good Practice (Richard Speed)
12. Modeling with LISREL: A Guide for the Uninitiated (A. Diamantopoulos)

Het boek is geschikt voor onderwijsdoeleinden, maar uitsluitend als introductie van de mogelijkheden van kwantitatieve technieken binnen marketing management. Per behandeld onderwerp bestaat er andere (nu al actuelere) literatuur die dieper op de technieken ingaat. De doelgroep bestaat dus volgens mij vooral uit marketing management studenten die weinig kwantitatieve achtergrond hebben. Ook kan dit boek wellicht gebruikt worden als tekst bij cursussen voor het bedrijfsleven, of MBA cursussen.

In relatie tot het veel gebruikte boek van Lilien, Kotler en Moorthy (Marketing models) is dit boek sterker gericht op (algemeen) econometrische technieken. Lilien et al besteedt ook aandacht aan beslissings-ondersteunende modellen, die in dit boek niet worden behandeld. Maar ook een aantal econometrische technieken zoals latente klasse modellen om consumenten heterogeniteit te modelleren worden helaas niet besproken. Het is goed om een tekst te hebben met uitleg van standaardtechnieken, maar door het ontbreken van een aantal onderwerpen is het gebruik van dit boek niet voldoende voor een goed overzicht van de kwantitatieve technieken in de marketing.

Verwijzingen:

- [1] Gary L. Lilien, Philip Kotler, K. Sridhar Moorthy, "Marketing Models", Englewood Cliffs, N.J. : Prentice-Hall, 1992.

*dr N. Piersma
Econometrisch Instituut
Erasmus Universiteit Rotterdam*

Johnson, R., Kuby, P.

Just the essentials of elementary statistics

1999, South Western, Belmont, California,

CD-ROM+590 blz, £ 25.99, ISBN 0-534-35779-2.

Robert Johnson is een bekende naam in de wereld van statistiek-onderwijs. In de eerste plaats als auteur, sinds kort co-auteur, van twee vernieuwende boeken in het grote Duxbury fonds: 'Elementary statistics', waarvan de achtste editie deze winter is verschenen, en het hier gerecenseerde 'Just the essentials of elementary statistics', afgelopen najaar op de markt gebracht. Daarnaast, of misschien kan ik beter zeggen: vooral, kreeg Johnson bekendheid als initiatiefnemer en organisator van de cyclus van jaarlijkse 'Beyond the Formula' conferenties, waarvan de vierde aflevering deze zomer plaats vindt. Zoals de naam van de conferentie al doet vermoeden, richt het zich op het onderwijs van de basiskennis statistiek aan niet-wiskunde studenten. Op een manier die zo min mogelijk de technieken en de bijbehorende formules benadrukt, maar zo veel mogelijk de statistiekstof conceptueel benadert. Om de richting nog wat verder aan te duiden: David Moore is iemand die zijn naam zeer nadrukkelijk aan de conferentie heeft verbonden.

Het leerboek ademt helemaal de geest van de conferenties. Een mooi geschreven tekst, met veel figuren, voorbeelden, opdrachten, toepassingen, etcetera. Een prachtig leerboek voor studenten die als een berg tegen statistiek opzien, en mede door overtuigende toepassingen moet worden gemotiveerd toch maar aan de slag te gaan (het zal duidelijk zijn: voor de meer formeel georiënteerde student is dit type leerboek een draak). In die zin is het leerboek zeer vergelijkbaar met de leerboeken waarvan David Moore co-auteur is.

De titel van het boek vind ik enigszins badinerend: het samenstel van woorden 'just', 'essential' en 'elementary' doet aan als een tautologie, die het basale aspect wel erg sterk benadrukt. Natuurlijk is het niet moeilijk te gissen waarom voor deze titel is gekozen: om het boek ten opzichte van het eerder geschreven boek duidelijk te positioneren, moest vooral het verschil belicht worden. Maar zo gek groot is dat verschil niet: beide boeken omvatten nagenoeg identieke eerste elf hoofdstukken, tezamen goed voor de stof van een traditionele introductie in de statistiek (beschrijvende statistiek, kansrekening, inferentiële statistiek tot en met chi-kwadraat toepassingen). 'Just the essentials of elementary statistics' houdt hier dan op, terwijl 'Elementary statistics' daar nog het enkelvoudige regressiemodel en een hoofdstuk niet-parametrische statistiek aan toevoegt.

Overigens is er nog een tweede verschil in beide boeken: ze gaan alle twee vergezeld van een 'multi-media learning resource' op CD-ROM, maar de elektronische hulpmiddelen die ze de lezer voorschotelen zijn verschillend. 'Elementary statistics' omvat een aan het leerboek aangepaste versie van CyberStats. Zelf heb ik enkel de β -versie van CyberStats gezien, maar die was zeer beloftevol: een volwaardig elektronisch leer- en oefenboek. In vergelijking daarmee viel de toevoeging bij 'Just the essentials of elementary statistics', die de naam StatSource draagt, wat tegen. Het

is enigszins een 'allegaartje', bestaande uit een aantal video's welke als introductie en motivatie voor het bestuderen van uiteenlopende statistische problemen kunnen worden gebruikt, data-bestanden, PowerPoint sheets, een deelverzameling van Gary McClelland's applets voor de simulatie van statistische concepten, etcetera. Leuk, maar niet zo aantrekkelijk als CyberStats. Terwijl het daarenboven nog wat technische onvolkomenheden bevat: het vereist de installatie van zowel QuickTime versie 3 als versie 2.1.2 naast elkaar (mist er één, dan lopen de video's niet, zonder enige waarschuwing) en laat af en toe je pc vast lopen.

In één opzicht is dit leerboek (samen met een heleboel andere) zeer Amerikaans: het is geschreven voor een inleidende cursus voor studenten afkomstig van verschillende 'schools'. Als gevolg zijn de voorbeelden en vraagstukken aan een breed spectrum van toepassingsgebieden ontleend: biologie, medicijnen, een beetje economie, ecologie,.... Voor een ieder wat wils, zou het credo kunnen zijn. Dat staat mijns inziens wat haaks op het uitgangspunt dat door voorstanders van de onderwijshervormingen nogal wordt benadrukt: ga in je onderwijs uit van de beleevingswereld van de studenten. In Nederland hebben de verschillende faculteiten meestal hun eigen cursussen inleidende statistiek. Dat staat ons toe om het principe 'aansluiten bij de belangstelling van de student' veel serieuzer gestalte te geven dan in dit soort boeken wordt gedaan: omdat alle studenten die de cursus lopen ofwel economie studeren (in mijn geval), of biologie, of ..., is het niet moeilijk die specifieke belangstelling te honoreren door voor een leerboek te kiezen dat voor die specifieke doelgroep is geschreven. En daar zijn er langzamerhand genoeg van. Juist door zo veel nadruk te leggen op de relevantie van het toepassingsgebied, ondergraven de hervormers de positie van dit type algemene leerboeken, hoe knap geschreven dan ook.

*D. Tempelaar
Capaciteitsgroep Kwantitatieve Economie
Universiteit Maastricht*

Huygens, C.

Van rekeningh in spelen van geluck

1998, Epsilon Uitgaven, Utrecht,
64 blz, fl 14.50, ISBN 90-5041-047-2.

Het is moeilijk om te besluiten wat er aan deze uitgave gerecenseerd dient te worden: het oorspronkelijk werk van Huygens, of het zeer te loven initiatief om dit werk voor de hedendaagse lezer bereikbaar en leesbaar te maken.

Ik zal beginnen met Huygens. Christiaan Huygens mag beschouwd worden als de grootste nederlandse wis- en natuurkundige, al bestond de qualificatie natuurkundige nog niet in zijn tijd. Bij het grote publiek is hij slechts bekend als 'de uitvinder van het slingeruurwerk'. Dit is onjuist, het slingeruurwerk bestond al, maar het was niet erg nauwkeurig. Dat werd veroorzaakt doordat de periode van de slinger afhangt van de uitslag. Huygens wist de link te leggen tussen dit probleem en het meer algemene

probleem van de isochrone curve, dat is een kromme met als eigenschap dat een kogel die over deze kromme naar beneden rolt daarvoor altijd dezelfde tijd nodig heeft, ongeacht het punt waarop de kogel wordt losgelaten. Dat was een probleem dat al door Galileï was bestudeerd. Huygens wist te bewijzen dat de cycloïde deze eigenschap heeft. Hij heeft toen het slingeruurwerk verbeterd door het bovenende van de starre slinger te vervangen door een leren riempje dat aan weerskanten afwikkelt over boogjes in de vorm van een cyloïde. Hierdoor volgt het uiteinde van de slinger een cycloïde baan en is daarmee isochroon.

Verder is hij bekend van het dispuut met Newton over het golf- versus het deeltjeskarakter van licht. Minder bekend is echter dat Huygens de eerste was die natuurkundige beginselen in de vorm van een formule beschreef, hij heeft daarbij ideeën ontwikkeld die pas algemeen bekend werden uit de publikaties van Newton, ca. 12 jaar later. Helaas was er een zeer koele, om niet te zeggen vijandelijke verstandhouding tussen beiden, men kan slechts fantaseren over wat er bereikt zou zijn als deze twee genieën hadden samengewerkt.

Ter gelegenheid van zijn 300-ste sterfdag werd in juli 1995 een twee-daags symposium georganiseerd waarin allerlei aspecten van zijn leven en zijn werk werden belicht. De voornaamste bijdragen hiertoe zijn gebundeld in een speciale uitgave van het tijdschrift 'de Zeventiende Eeuw'. Namens het International Statistical Institute dat bij deze herdenking was betrokken heb ik daar aandacht besteed aan het statistisch werk van Huygens en in het bijzonder aan 'Van Rekeningh in Spelen van Geluck'. Want ook op het gebied van de statistiek was Huygens een pionier! Hij heeft als eerste de formule $(a.p + b.q)/(p + q)$ voor een gewogen kans gepubliceerd. Net als in zijn natuurkundige werk was hij ook hier de eerste die een beginsel in een formule weergaf. Ik zal hier niet alles herhalen wat ik toen gezegd heb, maar mijn belangrijkste slotopmerking was dat *de tekst van Huygens zo helder en instructief is dat deze zonder meer aanbevolen kan worden voor gebruik in het hedendaagse onderwijs!* Met name het voortdurend terug grijpen op eerder bewezen zaken, iets dat we tegenwoordig recursie noemen, werd door mij geroemd.

Om dat qua instructie zo geniale gebruik van recursie te illustreren heb aan mijn zeer gemengde publiek, met hoofdzakelijk alpha's, het verhaal verteld van een test die aan een wiskundige en aan een natuurkundige wordt voorgelegd. De opstelling bestaat uit een aanrecht met waterkraan, gaskomfoor, een lege ketel en een doosje lucifers. De opdracht luidt: breng een ketel water aan de kook. Beiden vullen de ketel, zetten hem op het gas en steken dit aan. Daarna wordt de test herhaald, met als enige verschil dat de ketel nu gevuld is. De natuurkundige zet de gevulde ketel op het gas en steekt dit aan. De wiskundige giet slechts de ketel leeg en stelt dat hiermee de vorige situatie is verkregen die al was opgelost. Het was hierna niet moeilijk meer om mijn bewondering voor de manier waarop Huygens zijn stellingen bewees op mijn publiek over te brengen.

Ik zou kunnen volstaan met tevreden te constateren dat Wim Kleijne en uitgeverij Epsilon mijn wens hebben vervuld door de bijna 350 jaar oude tekst op een bijzonder fraaie wijze uit te geven en daarmee inderdaad bereikbaar te maken voor velen. Maar

daarmee doe ik Kleijne tekort. Hij heeft wezenlijk toegevoegd aan dat wat Huygens schreef. Hij heeft niet alleen de tekst in hedendaags nederlands omgezet (een facsimilé van het origineel is er naast afgedrukt), maar daarnaast door een aantal noten en bijlagen de bruikbaarheid nog meer vergroot dan deze al was.

Het boekje beveelt zichzelf aan voor o.a. zelfstandige bestudering in het VWO. In ben het daar van harte mee eens, maar ik zou graag zien dat het in een veel bredere kring bekend wordt. De lage prijs van deze uitgave laat een zeer ruime verspreiding toe. Het is zonder meer een uitstekend cadeau voor iedereen die ook maar een beetje in statistiek geïnteresseerd is. Statistisch Nederland kan niet trots genoeg zijn op haar pionier Huygens. Het zou daarom niet misplaatst zijn de hoofdlezing op jaarlijkse statistische dag nieuwe stijl de naam 'Huygens-lezing' te geven.

G.J. Stemerding

ISI, Senior Executive Corps

Aven, T., Jensen, U.

Stochastic models in reliability

1999, Applications of mathematics stochastic modelling and applied probability,
Springer-Verlag, Berlin,
blz, DM 119.00, ISBN 0-387-98633-2.

Via de VVS
25% korting!

Reliability and maintenance are studied in several mathematical subdisciplines, notably statistics, control theory, and stochastic processes. This book focuses on the latter. A solid foundation for reliability models is laid using renewal theory and martingales, and the power of these methods are shown for various reliability and maintenance models. The book concentrates on the application of the mathematical techniques to reliability models, and is thus mainly of interest to researchers in the fields of stochastic processes and reliability models. It is well written, including a good introduction and many intuitive arguments, explaining the sometimes involved mathematical results. It also contains appendices on probability theory, that can be used to refresh the memory of readers when it comes to martingales and renewal processes.

"Stochastic Models in Reliability" starts with an introduction explaining intuitively the basic model used in this book. It is based on a failure rate representation of the system lifetime, where the failure rate is a stochastic process given the available information. The information plays an important role: e.g., do we have information on the components (with thus a failure rate that is a function of the states of the components) or do we only know the availability of the whole system (leading to the traditional deterministic failure rate)?

Chapter 2 considers some basic subjects in reliability models. System availability as a function of component availability is discussed, as well as IFR/DFR distributions, including closure theorems for monotone systems. These subjects can also be found in other introductory text on reliability or stochastic processes.

Chapter 3 introduces the mathematical model in detail, and shows how it can be used to model various systems involving reliability and maintenance. Chapter 4 deals

with the availability analysis of these models. Chapter 5 includes the possibility of repair. Again various standard concepts (minimal repairs, block replacement, etc.) are discussed in the context of the model of Chapter 3. The book is concluded with appendices on probability theory and renewal theory.

*G.M. Koole
Faculteit der Exacte Wetenschappen
Vrije Universiteit Amsterdam*

Gilliam, D.S., Picci, G.

Dynamical systems, control, coding, computer vision. New trends, interfaces and interplay

1999, Progress in systems and control,
Birkhäuser, Basel,
500 blz, DM 198.00, ISBN 3-7643-6060-7.

Dit lijvige boek is de 25e publicatie in de reeks 'PSCT: Progress in Systems and Control Theory'. Het is een reeks waarin jaarlijks n tot drie titels verschijnen, telkens met geselecteerde bijdragen aan één van de vele specialistische congressen / workshops die het vakgebied rijk is. Het congres waarvoor dit boek als een soort proceedings dient is het 'Mathematical Theory of Networks and Systems Symposium (MTNS 98)' gehouden in Padua, Italië, juli 1998. Het boek bevat bewerkingen van de plenaire lezingen, invitatie- lezingen en minicursussen en minisymposia. De twee eraan voorafgaande publicaties bevatten bijdragen aan de 'Conference on Stochastic Differential Equations', Győr, Augustus 1996 en de 'Second International Workshop on Optimal Design and Control', Arlington, Virginia, September 1997. Daarmee zij aangegeven dat de bijeenkomsten steeds een nogal specifiek thema kennen, en de ermee corresponderende congresbundels geven dat aspect goed weer: alhoewel geschreven door verschillende auteurs, is de samenhang van de bijdragen meestal groot.

De voorliggende 25e uitgave is in enige mate een breuk op deze traditie. Zoals de titel al aangeeft, is de waaijer aan onderwerpen zeer breed. Dat geldt ook voor het type bijdragen: die lopen uiteen van puur theoretische bijdragen aan de algebra zoals Fuhrmann's 'On Canonical Wiener-Hopf Factorizations', tot besprekingen van zeer specifieke praktische probleemstellingen als beeldreconstructie uit reliëf- of schaduwafbeeldingen. Of, om nog een ander contrast te noemen, van overzichtswerken als Sontag's omvangrijke 'Nonlinear Feedback Stabilization Revisited', of 'Stabilization of Nonlinear Systems Using Output Feedback' van Isidori, tot een beschouwing over 'an-isotropische afvlakking van posteriori kansen in modellen van MRI-data'.

De bundel is snel uitgebracht: een congres laat juli 1998, het boek op de markt aan het eind van dat zelfde jaar. In zeker één opzicht is dit proces te snel verlopen: het voorwoord, een inhoudsopgave met onderverdeling in boekdelen (Networks, Systems and Control Theory; System Theory and Coding; Vision; Discrete-Events and Hybrid Systems), en een lijst van auteurs is als een los inlegkatern aan het boek toegevoegd,

dus kennelijk in eerste instantie niet nodig geacht of, waarschijnlijker, gewoon vergeten. Wat mij betreft had er nog wel een inlegkatern bij gemogen, dan achterin het boek: bij zo'n enorme diversiteit aan onderwerpen werkt een index toch wel erg handig.

Voor wie is dit boek bestemd? Gegeven nogmaals die grote diversiteit in velerlei opzicht, zal een individuele lezer niet gauw meer dan één bijdrage lezen, laat staan de hele bundel. Als gevolg zal het boek eerder z'n plaats verdienen in onze (universiteits)bibliotheken, om daar gericht te worden geraadpleegd op de afzonderlijke bijdragen, dan in de eigen boekenshap van onderzoeker of docent.

D. Tempelaar
Capaciteitsgroep Kwantitatieve Economie
Universiteit Maastricht

van den Broek, J., Kop, P.
Kattenaids en statistiek

1999, Zebra,
 Epsilon Uitgaven, Utrecht,
 58 blz., fl 14.75, ISBN 90-5041-050-2.

Behandelde onderwerpen Dit boekje geeft een uitleg van statistische begrippen aan de hand van een onderzoek naar kattenaids.

Ik vind hoofdstuk 1 (Een onderzoek) wel duidelijk. Er wordt even duidelijk uitgelegd wat voor soort onderzoek het is, en wat het doel is van dit onderzoek. Er worden een aantal dingen over het onderzoek verteld, maar ik mis daar dan toch dat stukje over de onderzoeksvraag. Dat had beter in hoofdstuk 1 dan in hoofdstuk 2 gepast.

Op hoofdstuk 2 (Populatie en steekproef) heb ik niets aan te merken, het was duidelijk en overzichtelijk. De begrippen worden goed uitgelegd, en de formule die er wordt gegeven is volgens mij ook duidelijk.

Hoofdstuk 3 (Eigenschappen van een steekproeffractie) is niet zo goed in elkaar gezet als hoofdstuk 2. Het is namelijk een stuk rommeliger. Dat komt omdat er een aantal formules dwars door de tekst heen stond. In hoofdstuk 2 was dat geen probleem omdat het daar maar om één formule ging. De formules werden overigens wel goed uitgelegd.

Wat ik persoonlijk ook zou veranderen is: de twee samenvattingen midden in een hoofdstuk samenvoegen tot een samenvatting na de vragen.

Hoofdstuk 4 (Het toetsen van een hypothese) is een behoorlijk overzichtelijk en duidelijk hoofdstuk. Het enige wat hier misschien ook bij had gehoord is even een korte samenvatting aan het eind verder is het prima. Alles wordt goed uitgelegd en beschreven.

Hoofdstuk 5 (Het vergelijken van twee populatiefracties) is een prima hoofdstuk ook hier is alles goed ingedeeld. Het enige wat ook hier mist is een samenvatting.

Hoofdstuk 6 (Het vergelijken van twee populatiegemiddelden) is het beste hoofdstuk uit het boek. Er ontbreekt niets: er staan een duidelijke uitleg en een schema in zodat je voor je ziet wat je moet doen. Zo zouden alle hoofdstukken moeten zijn.

Doelgroep Dit boekje is volgens mij geschikt voor kinderen vanaf ongeveer 14 jaar die in een 3 HAVO of 3 Atheneum zitten. Dan worden deze onderwerpen ook besproken maar de boeken gaan niet zo diep in op dit onderwerp. Ik denk dat dit boekje gebruikt kan worden voor kinderen die extra wiskunde willen leren of dit boekje zou bij zo'n gewoon schoolboek geleverd moeten worden. Het onderwerp statistiek wordt duidelijker uitgelegd dus dan is het boekje ook wel goed voor kinderen die wiskunde moeilijk vinden, omdat er toch wat extra uitleg in zit.

Sterke punten

- Een duidelijke uitleg van de begrippen aan de hand van katten met aids.
- De vergelijking tussen kattenaids en statistiek.
- De volgorde waarin de hoofdstukken staan en de manier waarop het boek opgebouwd is.

Zwakke punten

- Aan het eind van sommige hoofdstukken mist een samenvatting.
- Er mist een lijst met alle formules die er in het boek gebruikt zijn.

De rest van de zwakke punten heb ik al genoemd onder de kop behandelde onderwerpen.

*Sophie Koning
Piter Jelles Openbare Scholengemeenschap
Leeuwarden*

Boertjens, J.H.S., Kroes, J.J.

Internetgids gezondheid

1999, Academic Service, Schoonhoven,
CDROM+288 blz, fl 29.90, ISBN 90-395-0990-5.

Volgens de auteurs is het boek geschreven voor iedereen die met gezondheid te maken heeft dus eigenlijk voor iedereen. Deze opvatting heeft er echter toe geleid dat de indeling van het boek op sommige plaatsen wat rommelig is. De indeling van de hoofdstukken had strakker gekund. Zo gaat hoofdstuk 4 gaat over ziekte oftewel fysieke gezondheid en hoofdstuk 7 over geestelijke gezondheid. Deze twee hadden in

een hoofdstuk "gezondheid" met een onderverdeling naar fysiek en geestelijk kunnen worden onderverdeeld.

Het boek bestaat uit 10 hoofdstukken, 2 bijlagen en een CD-rom. De eerste drie hoofdstukken gaan over het Internet zelf en bespreken onder andere de eisen die aan een Internetcomputer worden gesteld, browsers, zoals Internet Explorer en Netscape, en zoekmachines. Behalve de informatie over zoekmachines zijn de eerdere hoofdstukken overbodig. Dit boek is namelijk bestemd voor mensen die informatie over gezondheid op het Internet willen vinden en dan zou ik ervan uitgaan dat ze de basis voor het benaderen van het Internet al onder de knie hebben.

De volgende hoofdstukken bespreken de internetsites voor verschillende onderwerpen. Hoofdstuk 4 gaat over ziekten, hoofdstuk 5 over mensen en gezondheid, hoofdstuk 6 over seksualiteit, zwangerschap en geboorte, hoofdstuk 7 over geestelijke gezondheidszorg, hoofdstuk 8 over alternatieve gezondheidszorg, hoofdstuk 9 over voorlichting bijzondere onderwerpen en hoofdstuk 10 gaat over voorlichting algemeen.

Zowel in het boek als op de CD-rom wordt een groot aantal sites gegeven. Het grote probleem is dat er geen onderscheid wordt gemaakt tussen de verschillende sites. De auteurs zouden de sites hebben kunnen onderscheiden naar aard van de site: zoals wetenschappelijk, belangenverenigingen, lotgenoten etc. Een dergelijke onderverdeling zou een grote hulp zijn bij het zoeken op het internet.

De mogelijkheid van het vinden van tegenstrijdige informatie over gezondheid wordt in het boek helaas niet aan de orde gebracht. Dit is een gemiste kans, omdat het vinden (krijgen) van tegenstrijdige informatie tot veel problemen kan leiden. Zo kan het vinden van tegenstrijdige informatie over overlevingskansen bij bepaalde aandoeningen grote gevolgen hebben (zie Moorer en Kamps). De auteurs zouden een korte handleiding voor het omgaan en evalueren van internetsites in een volgende editie kunnen overwegen.

Een aantal van de besproken internetsites zijn helaas met de opgegeven internetaadressen niet meer bereikbaar. Sommige zijn nog wel te achterhalen, maar een deel is echt niet meer te vinden. Dit probleem was gezien de veranderlijkheid van het internet niet te voorkomen.

Ondanks mijn kritische bespreking denk ik dat het een groot aantal mensen kan helpen om die informatie over gezondheid te vinden die ze zoeken. Het aantal meegeleverde links is bijzonder groot en als startpunt voor een zoektocht op het internet zullen ze ruim voldoende zijn.

Verwijzingen:

- [1] Moorer P. en Kamps W.A. (submitted) The effect of statistics on trust in medical professionals.

*P. Moorer
ARGO b.v.
Groningen*

Akaike, H., Kitagawa (eds)

The practice of time series analysis

1999, Statistics in engineering and physical science,

Springer-Verlag, Berlin,

400 blz, DM 99.00, ISBN 0-387-98658-8.

Via de VVS
25% korting!

From the Preface to the English Version: "This book is aimed at presenting examples of the successful application of time series analysis and control in various fields such as engineering, earth science, medical science, biology and economics." "The original version of this book was published in Japanese, in two volumes in 1994 and 1995, respectively."

The publication was planned as one of the activities to celebrate the 50th anniversary of the Institute of Statistical Mathematics in Tokyo. The book presents a broad overview of serious real life applications of likelihood based statistical time series analysis in Japan using methodologies which were developed at the Institute of Statistical Mathematics. The best known element of these methodologies is the Akaike information criterion in the context of order selection in Autoregressive models. Unfortunately many of the other modelling strategies are less widely known outside Japan since many interesting applications were published in Japanese. It was an excellent idea to publish an English edition of the book as well. It can be viewed as an up-to-date companion to the "basic general references", Akaike and Nakagawa (1988), Kitagawa and Gersch (1996), and Sakamoto, Ishiguro, and Kitagawa (1986), but basic knowledge of time series analysis is sufficient to understand most of the book. It ends with a short methodological chapter by Akaike and a short technical appendix by Kitagawa. It was published in a new Springer series "Statistics for Engineering and Physical Science", which already has been renamed to "Statistics for Engineering and Information Science". Note that recent publications on the book's topic are available via the Institute's website at <http://www.ism.ac.jp>, which has an extensive English section. Important software used by the authors of the book is available freely for academic use as well.

The applications in the book consider original empirical time series data sets in very different fields of science. Time series modelling is used to answer relevant research questions in earth science, biomedical science, engineering and econometrics. Most data sets are Japanese. The research is presented in articles of 10 to 20 pages.

Most applications use linear multivariate time series models with up to six state-variables. The multiple Autoregressive model (or Vector AR) is most frequently used, allowing for time varying parameters when necessary. The basic model selection is likelihood based, but appropriate Bayesian smoothness priors are used to identify processes with time varying parameters. State space models with Gaussian errors are used in many other applications. Many interesting results can be obtained using linear models, but nonlinear models are used when the data make this necessary. Survival data, data with outliers and data with stochastic switching require non-Gaussian distributions and other nonlinear state space representations. The more sophisticated nonlin-

ear models are confined to small-scale applications due to the existing lack of speed for the simulations and computations involved. Although part of the data analysis and model interpretation is done in the frequency domain, estimation is performed in the time domain in most cases.

The most important goals of the statistical analyses are the identification of impulse responses, transfer functions, (variance) decompositions and predictions. This involves a priori transformations of the data, the empirical identification of feedback mechanisms from closed-loop systems in engineering applications, order selection of AR models, the identification of the number of subintervals for locally stationary AR models, selecting explanatory variables, identification of the degree of smoothness of stochastic trend components and many other modelling choices that are more application specific.

Although automatical procedures have been derived for many of these tasks "it does not follow that it will automatically produce useful outcomes", as pointed out by Akaike in the conclusion of the book. It is stimulating to see so many useful outcomes for clear, purposeful applications in one volume. Practical criteria for identification dominate: only sensible decompositions, control systems, coherencies and predictions are retained. Theoretical selection criteria like A(B)IC often turn out to be helpful in the structure of the models concerned, but they are not the main part of the analysis.

Most authors used versions of specially developed FORTRAN computer programs like TIMSAC (Time Series Analysis and Control) BAYTAP (Bayesian Tidal Analysis Program). Most results are presented in sufficiently many clear figures.

The book contains 23 chapters by 23 different authors. It has over 400 pages. Each chapter is to-the-point. The collection of data and the research problems are described in sufficient detail, so that "laymen" from other disciplines can grasp the purpose of the model and understand the motivation and results, and maybe apply the ideas in their own area. It is one of the attractions of time series analysis that similar models have found successful applications in different disciplines.

The models are clearly presented in their mathematical form, often in state space, and in flow charts. The modelling strategy is always discussed. Computational details are omitted. For this purpose one has to consult the documentation of the main software suites and the basic references. Naturally all applications present (time varying) characterizations of the dynamics, impulse responses, decomposition of means and/or variances (so-called relative power contributions) of the variables under study.

The applications are divided across disciplines as follows. There are eight chapters on control and engineering. These concern the design of (power) plants and the empirical analysis of dynamic reactions of vehicles, bicycles and ships to external shocks. It includes a chapter on the control cement facturing, the research problem that led to development of the AIC. Five chapters concern medical science and biology. They vary from applications in neuroscience to pharmacokinetics. Six chapters relate to the earth Sciences. They study the dynamics of groundwater levels and earth quakes, tidal waves and satellite data. There are three econometric chapters, two on macroeconomic data and one on financial data. There is a map ordering the chapters according to model

and discipline in the beginning of the book. All chapters can be read independently.

Most applications form part of long-lasting research projects and are therefore far from trivial. The book contains few examples with simple data sets that could be replicated easily by a single researcher.

Although the English translation is quite adequate in general, it cannot hide that most chapters have not been written in English originally.

There are of course many other useful approaches to time series modeling, different from the ones used in the book. It is refreshing to read about so many up-to-date applications that do not even mention the ARIMA model, considered by many as a cornerstone of introductory time series courses.

The book should appeal to many researchers interested in reasonably advanced empirical time series modelling. I found the multi-disciplinary content stimulating.

Verwijzingen:

- [1] Akaike, H. and T. Nakagawa (1988). *Statistical Analysis and Control of Dynamic Systems*. Dordrecht: Kluwer.
 - [2] Kitagawa, G. and W. Gersch (1996). *Smoothness Priors Analysis of Time Series*. New York: Springer-Verlag.
- Sakamoto, Y., M. Ishiguro, and G. Kitagawa (1986). *Akaike Information Criterion Statistics*. Dordrecht: D. Reidel.

dr M. Ooms
Econometrisch Instituut
Erasmus Universiteit Rotterdam

Belsley, D.A.

Conditioning diagnostics; collinearity and weak data in regression

1991, John Wiley & Sons Ltd., Chichester,
xx+396 blz, £ 47.50, ISBN 0-471-52889-7.

The classical standard linear regression model is the starting point for the presentation of the issues the book addresses. Especially the use of this model in actual practice and its estimation by least squares is the main theme. The problems discussed are very familiar to econometricians and statisticians who do not have the possibility to design their experiments. Let $y = X\beta + \epsilon$ denote the standard linear regression model with the usual assumptions. In applications it frequently happens that X is observed and the researcher cannot influence the measurements in X . In case of design of experiments the design matrix X is explicitly constructed by the experimenter and problems like (multi)collinearity can be avoided. The problem of collinearity is often recognized from the occurrence of high standard errors, low t statistics, nonsensical or overly sensitive parameter estimates. It is difficult to have a good picture of what causes the collinearity. The use of the correlation matrix of the explanatory variables may show e.g. a strong relation between two of them, but quite frequently the pairwise correlations will be moderate and a satisfactory explanation cannot be given this way.

One sees the effects of collinearity, but not the collinearity itself. The book provides tools to make the picture complete.

Collinearity indicates the existence of a nearly linear relationship among a set of (explanatory) variables. The author shifts the attention to the concept of conditioning that encompasses the term collinearity. Conditioning is concerned with the sensitivity (or insensitivity) of a given relation to perturbations in the underlying data. This is also related to numerical problems that may occur in calculating the least squares estimates, because $X'X$ is e.g. nearly singular. The author more or less tacitly assumes that the model has a good specification: there are no (theoretical) arguments to delete variables. The author describes a procedure that gives the complete picture of the collinearity. The conditioning diagnostics are based on the matrix X that has been rescaled in such a way that the columns have lengths one. Mean centering is highly dissuaded! The author gives an extensive discussion on these points. A table that gives the so-called scaled condition indexes and variance-decomposition proportions can be used to look for degrading collinearity. To complete the picture, auxiliary regressions are used: the variables that are involved in near dependencies are regressed on the other variables. This procedure is extensively illustrated.

It depends on the error variance whether collinearity (whenever it is present) is harmful. In order to detect harmful collinearity the concept of signal-to-noise is used. Harmful collinearity occurs in case of inadequate signal-to-noise. Briefly formulated: the ratio of the least squares estimate of a coefficient to its standard deviation is too small. The terminology short data is used for inadequate signal-to-noise in the situation that there is no collinearity present according to the collinearity diagnostics.

The chapters 8 - 11 discuss various related topics. Chapter 8 treats the interplay between influential observations and collinearity. An influential observation can mask collinearity, but it can also create it. Chapter 9 is about collinearity diagnostics in models with logarithms and first differences. Chapter 10 discusses possible corrective actions. Case studies enrich the exposition. Chapter 11 generalizes to nonlinear and simultaneous-equations models. In this chapter he also addresses the question when to use LS, 2SLS or 3SLS.

The author has done his utmost to explain, to describe and to illustrate concepts like collinearity, weak data, conditioning diagnostics, etc. To this end he presents theory, artificial examples, simulations and real applications. The author's arguments are very convincing. In the preface the author rightly claims that the book is suitable for academic course work and useful as a research monograph. The book is invaluable for everyone who applies linear regression.

*prof. dr A.G.M. Steerneman
Vakgroep Econometrie
Rijksuniversiteit Groningen*

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.

Hunt, P.J., Kennedy, J.E.

Financial derivatives in theory and practice

2000, Wiley series in probability and statistics,
John Wiley & Sons Ltd., Chichester,
376 blz, £ 55.00, ISBN 0-471-96717-3.

German, R.

Performance analysis of communication systems. Modeling with non-Markovian stochastic Petri nets

2000, John Wiley & Sons Ltd., Chichester,
440 blz, £ 55.00, ISBN 0-471-49258-2.

Saltelli, A., Chan, K., Scott, E.M. (ed.)

Mathematical and statistical methods for sensitivity analysis

2000, John Wiley & Sons Ltd., Chichester,
504 blz, £ 55.00, ISBN 0-471-99892-3.

van de Craats, J.

Vectoren en matrices, een inleiding in de lineaire algebra

2000, Epsilon Uitgaven, Utrecht,
xvi+222 blz, fl 37.50, ISBN 90-5041-056-1.

Babbie, E.R., Halley, F., Zaino, J.

Adventures in social research. Data analysis using SPSS for Windows 95/98

200, Sage, Thousand Oaks, California,
480 blz, £ 27.00, ISBN 0-7619-8676-6.

Argyrous, G.

Statistics for social and health research. With a guide to SPSS

2000, Sage, Thousand Oaks, California,
560 blz, £ 24.99, ISBN 0-7619-6818-0.

Kanwal, R.

Singular integral equations

2000, Birkhäuser, Basel,
viii+427 blz, DM 138.00, ISBN 3-7643-4085-1.

Shumway, R.H., Stoffer, D.S.

Time series analysis and its applications

2000, Springer texts in statistics,
Springer-Verlag, Berlin,
xii+549 blz, DM 159.00, ISBN 0-387-98950-1.

Via de VVS
25% korting!

Bapat, R.B.

Linear algebra and linear models

2000, Springer-Verlag, Berlin,
x+138 blz, DM 89.00, ISBN 0-387-98871-8.

Via de VVS
25% korting!

Salkind, J.

Statistics for people who (think they) hate statistics

2000, Sage, Thousand Oaks, California,
408 blz, £ 19.99, ISBN 0-7619-1622-9.

Shorack, G.R.

Probability for statisticians

2000, Springer texts in statistics,
Springer-Verlag, Berlin,
xvii+585 blz, DM 159.00, ISBN 0-387-98953-6.

Via de VVS
25% korting!

Whittle, P.

Probability via expectation. 4th ed.

2000, Springer texts in statistics,
Springer-Verlag, Berlin,
xxi+352 blz, DM 129.00, ISBN 0-387-98955-2.

Via de VVS
25% korting!

Krause, A., Olson, M.

The basics of S and S-PLUS. 2nd ed.

2000, Statistics and computing,
Springer-Verlag, Berlin,
xxii+382 blz, DM 98.00, ISBN 0-387-98961-7.

Via de VVS
25% korting!

Dodge, Y., Jureckova, J.

Adaptive regression

2000, Springer-Verlag, Berlin,
xii+177 blz, DM 98.00, ISBN 0-387-98965-X.

Via de VVS
25% korting!

Pinheiro, J.C., Bates, D.M.

Mixed effects models in S and S-PLUS

2000, Statistics and computing,
Springer-Verlag, Berlin,
xvi+528 blz, DM 139.00, ISBN 0-387-98957-9.

Via de VVS
25% korting!

Bhatti, M.A.

Practical optimization methods with Mathematica applications

2000, Springer-Verlag, Berlin,
xii+715 blz, DM 159.00, ISBN 0-387-98631-6.

Via de VVS
25% korting!

Venables, W.N., Ripley, B.D.

S Programming

2000, Statistics and computing,
Springer-Verlag, Berlin,
x+264 blz, DM 129.00, ISBN 0-387-98966-8.

Via de VVS
25% korting!

Snijders, T.A.B., Bosker, R.J.

Multilevel analysis. An introduction to basic and advanced multilevel modeling

1999, Sage, Thousand Oaks, California,

272 blz, £ 18.99, ISBN 0-7619-5890-8.

Field, A.

Discovering statistics using SPSS for Windows. Advanced techniques for beginners

2000, Sage, Thousand Oaks, California,

512 blz, £ 19.99, ISBN 0-7619-5755-3.

Taris, T.W.

A primer in longitudinal data analysis

2000, Sage, Thousand Oaks, California,

176 blz, £ 18.99, ISBN 0-7619-6027-9.

Tijms, H., Heierman, F., Nobel, R.

Poisson, de Pruisen en de lotto

2000, Zebra-recks,

Epsilon Uitgaven, Utrecht,

54 blz, fl 16.75, ISBN 90-5041-059-6.

Rossman, A., Chance, B.L.

Workshop statistics: discovery with data

2000, Springer-Verlag, Berlin,

625 blz, DM 79.00, ISBN 1-930190-03-4.

Via de VVS
25% korting!

Nolan, D., Speed, T.P.

Stat labs. Mathematical statistics through applications

2000, Springer texts in statistics,

Springer-Verlag, Berlin,

viii+282 blz, DM 69.00, ISBN 0-387-98974-9.

Via de VVS
25% korting!

Edwards, D.**Introduction to graphical modelling**

2000, Springer texts in statistics,
 Springer-Verlag, Berlin,
 xv+333 blz, DM 139.00, ISBN 0-387-95054-0.

Via de VVS
 25% korting!

Pelsser, A.**Efficient methods for valuing interest rate derivatives**

2000, Springer finance,
 Springer-Verlag, Berlin,
 blz, DM 129.00, ISBN 1-85233-304-9.

Via de VVS
 25% korting!

Behling, O., Law, K.S.**Translating questionnaires and other research instruments**

2000, Quantitative applications in the social sciences,
 Sage, Thousand Oaks, California,
 80 blz, £ 9.99, ISBN 0-7619-1824-8.

Fox, J.**Multiple and generalized nonparametric regression**

2000, Quantitative applications in the social sciences,
 Sage, Thousand Oaks, California,
 96 blz, £ 9.99, ISBN 0-7619-2189-3.

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/>, knop *Ledenservice*.

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 zinige 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)	

