

Boekbesprekingen

Redactie

Alex J. Koning

postadres: Vakgroep Econometrie en Besliskunde
Erasmus Universiteit Rotterdam
Postbus 1738
3000 DR Rotterdam
telefoon: 010-4081268/59
fax: 010-4527746
e-mail: koning@few.eur.nl

Besproken boeken in Kwantitatieve Methoden nr. 55

Saville, D.J., Wood, G.R.

Statistical methods. A geometric primer

Simonoff, J.S.

Smoothing methods in statistics

Stapleton, J.H.

Linear statistical models

Horst, R., Tuy, H.

Global optimization. Deterministic approaches (3rd ed.)

Lee, M.-J.

Method of moments and semiparametric econometrics for limited dependent variable models

Mittelhammer, R.C.

Mathematical statistics for economics and business

Carroll, R.J., Ruppert, D., Stefanski, L.A.

Measurement error in nonlinear models

Davidian, M., Giltinan, D.M.

Nonlinear models for repeated measurement data

Waterman, M.S.

Introduction to computational biology

Guttorp, P.

Stochastic modelling of scientific data

Shiryaev, A.N.

Probability

Rossman, A.J.

Workshop Statistics

Valkenburg, H.A., Hofman, A.

Een kwart eeuw Hippocratische epidemiologie

Longford, N.T.

Models for uncertainty in educational testing

Roes, C.B.

Shewhart-type charts in statistical process control

Fisher, L.D., van Belle, G.

Biostatistics. A methodology for the health sciences

Saville, D.J., Wood, G.R.

Statistical methods. A geometric primer

1996, Springer-Verlag, Berlin,

275 blz, DM 68.00, ISBN 0-387-94705-1.

Het doel van dit boek is een zo eenvoudig mogelijke presentatie van de wiskunde (lees: geometrie) waarop elementaire statistische methoden, gerelateerd aan de normale verdeling, gebaseerd zijn. Het boek is zowel geschikt voor studenten statistiek, ter aanvulling van meer gebruikelijke boeken wanneer men echt een goed begrip voor de materie wil krijgen, als voor studenten in de lineaire algebra of geometrie, die deze statistische methoden als aardige toepassingen van hun studie onderwerpen kunnen beschouwen. Het boek kan beschouwd worden als een zustertekst van het boek 'Statistical Methods: The Geometric Approach' (1991, zelfde auteurs en uitgever), waarbij het oudere zusje als 'wijzer' en 'meer volwassen' beschouwd kan worden.

Het huidige boek is vrij beperkt qua onderwerpen, het telt naast 'Introduction' en 'Overview' vier goede hoofdstukken: 'Paired Samples', 'Independent Samples', 'Several Independent Samples (Analysis of Variance)' en 'Simple Regression'. Ieder van deze hoofdstukken geeft inzicht in de geometrie die aan de basis van de gebruikelijke statistische methode ligt, met nadruk op het testen van hypothesen (inderdaad, t en F tests). Per hoofdstuk wordt een data-set geïntroduceerd, kort besproken (bv. voor welke populatie de data-set representatief zou kunnen zijn), en dan stapsgewijs geanalyseerd. In wezen is de gehele gebruikte geometrie niet meer of minder dan 'Pythagoras', maar het boek zal voor menig student een nuttig hulpmiddel zijn bij een inleidende cursus statistiek. Ook de appendices zijn nuttig, met name voor studenten die hun voorkennis (App. A: 'Geometric Tool Kit', B: 'Statistical Tool Kit') weer eens moeten oprakelen, en voor studenten die nog niet vertrouwd zijn met de nodige methoden voor 'Computing' (App. C). Er wordt een aardige alternatieve test methode voorgesteld in appendix D, velen zullen dit over willen slaan.

Voor docenten heeft het boek overigens ook heel wat te bieden. Niet alleen is de geometrische aanpak aantrekkelijk voor het presenteren van deze basismethoden, het boek heeft veel aardige opgaven, waarvan een deel met nauwkeurig uitgewerkte oplossingen, en ieder hoofdstuk eindigt met een 'Class Exercise'. Door deze oefening inderdaad klassikaal te laten plaatsvinden, ontstaat een mooie overbrugging tussen het presenteren van de stof door de docent en het maken van de opgaven door de studenten afzonderlijk. Zelfs de appendices A, C en D zijn voorzien van opgaven.

Het boek leest heel prettig, ziet er aardig uit (soft-cover, Latex lay-out, het heeft wel iets van een proefschrift), en is een welkome aanvulling in elke wiskunde/statistiek bibliotheek. Er is één grote 'maar': het boekje is veel te duur, zowel voor wat het biedt als voor de studenten die de hoofdzakelijke doelgroep vormen. Voor één-derde van de prijs zou ik het studenten aanbevelen als een nuttige aanschaf, voor het equivalent van

DM 68.00 (of meer) raad ik hen aan het eens een weekje uit de bibliotheek te lenen.

*dr F.P.A Coolen
Department of Mathematical Sciences
University of Durham*

Simonoff, J.S.

Smoothing methods in statistics

1996, Springer series in statistics,
Springer-Verlag, Berlin,
355 blz, DM 84.00, ISBN 0-387-94716-7.

Dit boek behandelt 'smoothing' (gladstrijkings) methoden in de statistiek. Het is onderverdeeld in 7 hoofdstukken en een uitgebreide literatuur lijst van maar liefst 30 bladzijden. De onderwerpen variëren van univariate dichtheid schatting (eenvoudig en gladgestreken), multivariate dichtheid schatting, verdelingsvrije regressie tot gladstrijking van categorische gegevens. In het laatste hoofdstuk worden nog enkele meer specialistische onderwerpen kort behandeld, o.a. de 'gladgestreken' bootstrap.

Het boek volgt een vaste structuur. Ieder hoofdstuk is in 3 onderdelen verdeeld: presentatie van het onderwerp, achtergrond materiaal en opgaven. Het eerste onderdeel bevat een kort overzicht van de procedures waarbij verschillende varianten aan de orde komen. Monte Carlo technieken worden gebruikt, bijvoorbeeld om betrouwbaarheids intervallen vast te stellen. Het tweede onderdeel is vooral bedoeld als toegang tot de literatuur. Soms wordt nog even ingegaan op de onderwerpen zelf. De opgaven in het derde onderdeel zijn aardig, en kunnen zonder al te veel problemen gemaakt worden, althans bij aanwezigheid van de juiste software waarvan achterin in het boek een overzicht wordt gegeven.

De kracht van het boek zit hem ongetwijfeld in het gemak en de grote leesbaarheid waarmee de technieken uitgelegd worden, maar vooral ook in het grote aantal voorbeelden. Deze stammen uit de meest uiteenlopende onderwerpen. Van de 'Old Faithful' geiser in Yellowstone, via de kwaliteit van basketbal spelers en vervalste Zwitserse bankbiljetten tot diabetes patienten. En het aardige is dat deze gegevens gewoon via het net beschikbaar zijn. Op de home-pagina van de schrijver vinden we de laatste errata, aanvullingen (en correcties) op de literatuur lijst en op de computer berekeningen. De bestanden zijn eenvoudig op te roepen en de analyses zijn snel te verifiëren.

De schrijvers hebben zich vooral moeite getroost de methoden te presenteren, maar niet zozeer om ze uit te leggen. Dit heeft geleid tot een kloof tussen voorbeelden en theorie. De werking van de technieken wordt beschreven en met gegevens geïllustreerd. Soms wordt dan verrassend geconcludeerd dat we 'nu echt wat zien' wat we 'anders vast niet gezien zouden hebben'. Dat klopt dan ook vaak, waarmee m.i. de belangrijkste reden om smoothing technieken te gebruiken geformuleerd kan worden: ze zijn bijzonder geschikt om gegevens grafisch samen te vatten en enkele belangrijke

kenmerkende eruit te lichten. Asymptotische beschouwingen, die tot een minimum beperkt zijn gebleven, tonen aan dat de methoden in ieder geval asymptotisch consistent zijn. De sprong tussen het etaleren van de technieken en de asymptotiek is daarmee groot. De asymptotiek is lang niet altijd even eenvoudig te volgen, wordt eigenlijk nooit bewezen, en is daarmee vaak niet goed te begrijpen, laat staan toe te passen op eigen gegevens.

Ik had natuurlijk de stille hoop dat dit boek een aardige beschouwing zou bevatten over kriging in relatie tot smoothing technieken, aansluitend op een recente studiedag in Wageningen. In het boek 'generalized additive models' van Hastie en Tibshirani (Chapman and Hall, London) wordt dit kort vermeld zonder dat het daar uitgewerkt wordt. In dit boek komt het helemaal niet aan de orde. Twee-dimensionale kernels worden wel behandeld, maar de relatie naar kriging wordt verder niet gelegd.

Al met al vormt dit boek een nuttige ingang tot een interessante groep statistische methoden. Een enkele typfout (maar niet veel) en af en toe een slordigheidje in de plaatjes doen aan de kwaliteit van het boek niet veel af. Zeker voor hen die in de praktijk met smoothing technieken te maken krijgen is het goed om het binnen handbereik te hebben: de technieken staan best beschreven, de asymptotische rechtvaardiging kun je zonder probleem overslaan en de geïnteresseerde belangstellende is met een paar simpele opmerkingen de bibliotheek in te sturen om 'het nog maar eens voor jezelf na te lezen'. De keuzes van smoothing parameters worden met praktijkervaringen ondersteund. Voor hen die voor het eerst met de methoden te maken krijgen is het zeker aardig om verbanden met andere studies te zien. Het boek vormt daarmee een prima ingang tot een verder begrip en een goede verspreiding van deze nieuwe, populaire methoden.

*prof. dr A. Stein
Landbouwwuniversiteit Wageningen*

Stapleton, J.H.

Linear statistical models

1996, Wiley series in probability and statistics,
John Wiley & Sons Ltd., Chichester,
xiii+449 blz, £ 50.00, ISBN 0-471-57150-4.

Dit leerboek uit de reeks "Probability and Statistics" van Wiley kent een achttal hoofdstukken, waarvan de inhoud als volgt kan worden omschreven:

H1 Linear Algebra, Projections: introductie lineaire algebra, projecties;

H2 Random Vectors: introductie waarschijnlijkheidsrekening;

H3 The Linear Model: na een bespreking van de lineaire hypothese, met daarin opgenomen een 40-regelige filosofie voor statistici, volgen betrouwbaarheidsintervallen en toetsen, het Gauss-Markov theorema, de (geometrische) interpretatie van regressiecoëfficiënten en de correlatiecoëfficiënt.

- H4 Fitting of Regression Models: hierin aandacht voor onderwerpen uit de modelselectie en alternatieve regressietechnieken, zoals: linearisatie, residu-analyse, col-lineariteit, niet-lineaire kleinste kwadraten, robuuste regressie, bootstrapping.
- H5 Simultaneous Confidence Intervals: vergelijking van Bonferroni, Scheffe, Tukey en Bechhofer's methoden.
- H6 Two-Way and Three-Way Analysis of Variance: twee- en drie-weg variantie-analyse, inclusief de situatie met ongelijk aantal waarnemingen per cel.
- H7 Miscellaneous Other Models: verbijzondert het eerder behandelde naar de volgende opzetten: 'random-effect model', 'nesting', 'split plot' en 'balanced incomplete block'.
- H8 Analysis of Frequency Data: met ondermeer het log-lineaire model en logistische regressie.

Een bespreking van dit boek van de hand van W.K. Wong verscheen in het juni-nummer van JASA (Wong, 1996). De recensent noemde daar als sterke punten van het boek: het gebruik van empirische problemen als illustratiemateriaal, de volledigheid van de behandeling (in het Engels zo mooi met self-contained aan te geven) en de opname van huiswerk-problemen en, deels, de oplossingen daarvan. Stapleton zelf voegt daar op de achterflap van het boek aan toe: de geometrische benadering, en het gebruik van de computerapplicaties SAS en S-Plus. Op ieder van deze veronderstelde sterke punten valt echter het nodige af te dingen. Het aantal voorbeelden is beperkt en daarenboven steeds ontleend aan andere bronnen, waaronder vooral de SAS-handleiding. Het 'zelfdragende' karakter van het boek betreft bij uitstek de benodigde voorkennis op het gebied der lineaire algebra en stochastiek (zie hoofdstukken 1 en 2), doch veel minder een aantal meer geavanceerde onderwerpen als principale componenten, canonische analyse, experimentele opzetten, Bayesiaanse aanpak, etc.. Op dat laatste vlak is een boek als Christensen (1996) zeker vollediger, doch dit ontbeert weer enig voorbeeld. Vraagstukken bevat het boek in ruime mate; de kanttekening is hier dat ze hoofdzakelijk rekenkundig van karakter zijn. Het geometrische element is wederom betrekkelijk bescheiden: het wordt af en toe benut voor het geven van een interpretatie, maar speelt nergens zo'n dominante rol als bijvoorbeeld in Saville en Wood (1991) (gerecenseerd in *Kwantitatieve Methoden* 39). En tenslotte: ook de rol van SAS en S-plus is zeer bescheiden. Als zwakker punt van het boek noemt Wong een aantal zetfouten en inconsistente notaties. Het valt niet moeilijk zijn lijstje van voorbeelden verder uit te breiden; anderzijds zijn ze nergens zeer storend en, voor een eerste druk, nogal onvermijdelijk.

Daar waar bovenstaande opsomming de indruk wekt dat de balans overslaat naar het negatieve, dient die conclusie met spoed de kop te worden ingedrukt: het leerboek van Stapleton is zeker een goede kandidaat als tekstboek voor een vak als Multipelle Regressie of Variantie-Analyse in de doctorale fase van de studie wiskunde, econometrie of een der natuurwetenschappen. Het onderscheidt zich van een aantal klassieke

teksten, en van de eerder genoemde Christensen (1987), door de aandacht voor het motiveren van de lezer. In tegenstelling tot Saville en Wood (1991) behoudt het daarbij een voldoende niveau van diepgang. Kan men zich vinden in de gekozen onderwerpen, dan verdient dit boek zeker een serieuze overweging. Want zo veel leerboeken op dit terrein die goed leesbare theoretische verhandelingen koppelen aan motiverende voorbeelden en inzicht verschaffende geometrische interpretaties zijn er zeker niet.

Verwijzingen:

- [1] Ronald Christensen: *Plane Answers to Complex Questions, The Theory of Linear Models*, 2nd ed. Springer, 1996.
- [2] David J. Saville, Graham R. Wood: *Statistical Methods: the Geometric Approach*. Springer, 1991.
- [3] Wen Kee Wong: Book Review 'Linear Statistical Models'. *JASA*, June 1996, p. 909-910.

*D. Tempelaar
Vakgroep Kwantitatieve Economie
Rijksuniversiteit Limburg*

Horst, R., Tuy, H.

Global optimization. Deterministic approaches (3rd ed.)

1996, Springer-Verlag, Berlin,
xviii+728 blz, DM 248,00, ISBN 3-540-61038-3.

Readership: students and researchers in numerical analysis, computer sciences.

In many practical problems scientists are confronted with standard global optimization problems of the following form. Given a nonempty closed subset D of the n -dimensional Euclidian space and a real continuous function $f : A \rightarrow \mathbb{R}$, where A is a suitable set containing D , find at least one point x^0 in D satisfying $f(x^0) \leq f(x)$ for all x in D , or show that such a point does not exist. If such a point exist, x^0 is called a global minimizer of the function f over the region D .

The literature offers many results that help to conclude that such a point x^0 exists. Often the conditions are related to continuity of f and rely on convexity/concavity properties of the function f and the feasible set D .

The literature about how to find such points x^0 is spread around and practitioners often have difficulties in finding a suitable algorithm for the problem at hand. This illustrates the main objective of this book. "... The enormous practical need for solving global optimization problems and the rapidly advancing computer technology has allowed one to consider problems which a few years ago would have been considered computationally intractable ... The goal of the book is to systematically clarify and unify the diverse approaches ..." The authors successfully obtain this goal.

The book starts with a survey of several large classes of global optimization problems and conclude that in many -not all- practical cases the optimization problem can be reduced to minimizing a concave function over a convex set. A lot of attention goes to the class d.c. of functions. This is the class of functions that can be expressed as the difference of two convex functions.

In order to solve optimisation problems, in general there are three types of approaches available. Let us denote by $P = (f, D)$ the problem of optimizing f over D and let $f^0(P)$ denote an optimal value of f .

Outer approximation methods are based on the construction of a sequence of problems $\{P(k)\}$ and values $f^0(P(k))$, where $P(k) = (f, D(k))$, where $D(k)$ is a decreasing sequence of subsets of D and converging to D and where $f^0(P(k))$ converges to $f^0(P)$. Outer approximation methods can be constructed by using different types of cutting procedures.

Dually, in inner approximation methods one approximates the feasible set D from the inside by a sequence of expanding sets $D(k)$ converging to D .

The branch and bound method is based on the following idea. We start from a 'nice' set $M(0)$ containing D and split ("branch") the set $M(0)$ into finitely many subsets $M(i)$. Now we consider the problem $P(i) = (f, D(i))$ where $D(i)$ is the intersection of $M(i)$ and D . We proceed by finding for each i an upper and lower bound for the function f . This way we can construct overall bounds $a \leq f^0(P) \leq b$. If $b - a$ is close to zero, we stop the algorithm. Otherwise we select one or more of the subsets $M(i)$ to obtain a refined partition of $M(0)$.

In the book the three types of approaches are used to give different algorithms to solve a wide variety of problems. The authors consider large scale problems, networks, bilinear problems, d.c.-programming problems, etc.

The technical prerequisites for the book are rather modest and are within reach for students with a sound knowledge of real analysis, linear algebra and convexity theory. The book is written well and by using a good combination of intuition and hard proofs, the result is a didactically attractive book. Each chapter begins with a summary of its contents and on several occasions the authors give or refer to suitable applications. The list of references is impressive and contains over 500 entries. The index list is modest but useful.

A minor point of regret is the absence of a list of algorithms discussed in the book. Also, with student-readers in mind, there are no exercises included. The value of the book could increase if each chapter ended with a section devoted to bibliographical notes. Also, as we live in a computer driven world, the book should give some guidelines and remarks about computer packages or programming languages. Finally, the book is devoted to deterministic approaches. In the next edition of the book there should be some place for the powerful stochastic approaches in global optimization.

prof. dr E. Omey
EHSAL

Lee, M.-J.

Method of moments and semiparametric econometrics for limited dependent variable models

1996, Springer-Verlag, Berlin,

xii+279 blz, DM 78,00, ISBN 0-387-94626-8.

Eén van de belangrijkste ontwikkelingen in de econometrie gedurende het einde van de jaren zeventig en het begin van de jaren tachtig is ongetwijfeld de ontwikkeling van modellen voor discrete en beperkt afhankelijke variabelen geweest. 'Standaard' regressiemodellen gaan er meestal van uit dat de te verklaren variabele continu is en dat maakt deze modellen minder goed toepasbaar als de keuze tussen een huur- of koopwoning (een 0/1-beslissing) of het arbeidsaanbod (een niet-negatieve variabele) moet worden gemodelleerd. In de loop van de jaren tachtig zijn nieuwe methoden ontwikkeld om dergelijke modellen te schatten, zonder stringente functionele vorm of verdelingsveronderstellingen te maken. In deze zin is dit boek van Lee de opvolger van het boek van Maddala (1983), waar de analyse van bijna alle modellen is gestoeld op de aanname van normaliteit van een storingsterm.

Het boek valt in twee delen uiteen. In het eerste deel wordt de gegeneraliseerde momenten methode behandeld, in het tweede deel komen semi-(non)parametrische methoden aan bod. In beide delen wordt begonnen met een discussie van de bouwstenen die nodig zijn om de schatters in latere hoofdstukken te begrijpen. In het eerste gedeelte wordt ook aandacht besteed aan de methode van de gesimuleerde momenten en varianten daarop. Het tweede gedeelte begint met niet-parametrische dichtheids-schatting en regressie. Vervolgens worden een groot aantal methoden behandeld om LDV-modellen te schatten zonder stringente veronderstellingen te maken. Aan bod komen onder meer: index-schatters, niet-parametrische instrumentele variabele schatters, rang-correlatie schatters, en nog vele anderen. Tenslotte worden in een appendix een aantal voorbeeldprogramma's gegeven om semi-parametrische schatters uit te rekenen, alsmede een kleine dataset. De programma's zijn geschreven in de matrixprogrammeertaal GAUSS. Helaas heb ik niet kunnen ontdekken of de programmacode en de dataset ergens op internet is gepubliceerd; met name het overtypen van vijf bladzijden data lijkt mij een onaangename en zinloze bezigheid.

Het boek is volgens de auteur bedoeld voor onderzoekers en eerste jaars AiO's. Beide groepen zullen dit boek niet als eerste tekst op dit gebied kunnen lezen, daarvoor is het te moeilijk. Enige voorkennis van ML-schatten van discrete keuzemodellen, asymptotiek, en kernschatters is eigenlijk wel vereist, hoewel de auteur zijn best doet om alle noodzakelijke technische bagage te behandelen. Het boek is zeer geschikt voor AiO's die eerder een cursus econometrie voor gevorderden hebben gevolgd.

Het boek heeft twee sterke punten en drie zwakke punten. Ten eerste staat er erg veel in het boek, de auteur geeft een breed overzicht van moderne schattingstechnieken. Dit is tevens een zwak punt: in zijn streven om zoveel mogelijk methoden te behandelen komen sommige methoden er zo bekaaid af dat een lezer die dat voor het eerst leest er weinig van zal snappen (de semi-nonparametrische schattingsmethode

van Gallant en Nychka wordt in zes regels en drie formules behandeld). Een ander goed aspect aan het boek is dat de auteur bewijzen van stellingen en beweringen niet schuwt. Dit heeft als voordeel dat de lezer beter in staat is om nieuwe artikelen in tijdschriften zelfstandig te lezen. Helaas heeft de auteur weinig kaas gegeten van toepassingen, daar waar hij spreekt over economische toepassingen (en dat is niet al te vaak) komt dat nogal geforceerd over. Tenslotte is de formuledichtheid in sommige paragrafen wel erg hoog, iets meer verbale toelichting zou geen kwaad hebben gekund.

Dit boek is een aanwinst voor de vakliteratuur omdat niet-triviale methoden op een toegankelijke wijze worden behandeld. Het boek kan worden beschouwd als een moderne tegenhanger van het boek van Maddala. Gelukkig heeft Lee zijn drukproeven beter bekeken zodat er weinig fouten in het boek staan. AiO's en onderzoekers op het gebied van de toegepaste microeconomie zullen veel plezier aan dit boek kunnen beleven omdat dit boek bij mijn weten als eerste een redelijke coherent en compleet overzicht geeft van semi-(non)parametrische microeconometrie.

Verwijzingen:

- [1] Maddala, G.S. (1983) *Limited-Dependent and Qualitative Variables in Econometrics* Cambridge (UK): Cambridge University Press.

dr R.H. Koning
Vakgroep Econometrie
Vrije Universiteit Amsterdam

Mittelhammer, R.C.

Mathematical statistics for economics and business

1996, Springer-Verlag, Berlin,
 xviii + 723 blz, DM 78,00, ISBN 0-387-94587-3.

Zoals de titel aangeeft is dit een boek waarin kansrekening en statistiek behandeld wordt voor een specifieke doelgroep die door de schrijver gekarakteriseerd wordt als studenten economie, econometrie en business (finance, marketing, accounting e.d.). De inhoudsopgave is als volgt:

1. Elements of probability theory.
2. Random variables, densities and cumulative distribution functions.
3. Mathematical expectation and moments.
4. Parametric families of density functions.
5. Basic asymptotics.
6. Sampling, sample moments, sampling distributions and simulation (trekken met en zonder terugleggen, empirische verdelingsfunctie, order statistics).

7. Elements of point estimation theory (sufficiency, completeness, Cramer-Rao ongelijkheid e.d.).
8. Point estimation methods (lineair model en kk schatters, meest aannemelijke schatters en momenten schatters).
9. Elements of hypotheses testing theory (Neyman-Pearson Lemma, UMP toetsen e.d.).
10. Hypothesis testing methods (GLR toets, Lagrange multiplier toets, Wald, betrouwbaarheidsintervallen, niet centrale χ^2 , t en F- verdeling).

Als we dit boek vergelijken met andere boeken die zich richten op een middenniveau en geen gebruik maken van maattheorie (ongeveer vergelijkbaar met het klassieke boek van Mood en Graybill), valt op dat er in dit boek meer details te vinden zijn dan gebruikelijk is in soortgelijke boeken van dit niveau.

Een voorbeeld hiervan is hoofdstuk 5 waarin naast diverse varianten van de zwakke wet van de grote aantallen en de centrale limietstelling ook convergentie met kans 1 en de sterke wet besproken wordt. Dit hoofdstuk begint met een paragraaf waarin Landau's O en o notatie nog eens uit de doeken gedaan wordt waarna de probabilistische equivalenten aan de orde komen. Als notatie voor convergentie in kans (a.s.) wordt zowel $\xrightarrow{P}$ ($\xrightarrow{a.s.}$) als plim (aslim) gebruikt, zoals in de econometrie gebruikelijk is.

In het kader van de toetsingstheorie komen wat ongebruikelijke zaken aan het licht, zoals de behandeling van een onderwerp als de Lagrange multiplier test. Een onderwerp als de tekentoets heb ik daarentegen niet kunnen vinden. Anderzijds worden andere verdelingsvrije toetsen zoals χ^2 toetsen, Kolmogorov-Smirnov, Shapiro-Wilks en Wald-Wolfowitz weer wel behandeld (in een deelparagraaf van hoofdstuk 10).

Het boek is goed leesbaar en nieuwe begrippen worden goed uitgelegd. De layout is overzichtelijk met tekstboxen voor de definities en aan het eind van ieder hoofdstuk een lijst van trefwoorden en vraagstukken (zij het zonder antwoorden, hetgeen studenten niet zo appreciëren) die veelal ingebed zijn in een vaak economische context. Voor meer ingewikkelde bewijzen wordt steeds een passende literatuurverwijzing gegeven.

Het boek bevat geen opgaven waarbij statistische pakketten nodig zijn voor de oplossing.

Door het gebruik van een dikke papiersoort oogt het boek nogal massief en ook de afmetingen doen aan een Amerikaans calculusboek denken.

Samenvattend is mijn oordeel positief voor de doelgroep van pure econometriestudenten die zich niet laten afschrikken door een formuleetje. Voorwaarde voor gebruik als collegemateriaal is wel dat er een behoorlijke hoeveelheid tijd nodig is om alles uit te leggen en te verteren. Voor andere studierichtingen met een hoog wiskundig gehalte is het boek wellicht wat eenzijdig, maar gezien de schaarste aan boeken op dit niveau

zeker de moeite waard om te bekijken.

dr J.L. Geluk

Vakgroep Econometrie en Besliskunde

Erasmus Universiteit Rotterdam

Carroll, R.J., Ruppert, D., Stefanski, L.A.

Measurement error in nonlinear models

1995, Monographs on statistics and applied probability 63,

Chapman & Hall, London,

305 blz., £ 29.95, ISBN 0-412-04721-7.

Measurement error in explanatory variables is a topic with a long history. It is amply discussed in the books of Kendall and Stuart. A classic in the context of linear regression is the book of W.A.Fuller, *Measurement error models*, Wiley 1987. Application in non-linear regression is of more recent date. Here, non-linear models should be understood as Generalized Linear Models, such as logistic regression, Poisson regression and survival analysis. The focus in this book is mainly on logistic regression. An important application is nutritional epidemiology, where the effect of dietary intake on prevalence and incidence of diseases is studied. It is practically impossible to obtain error-free data on intake of all food components. This has been a major problem for nutritional epidemiologists.

In the notation of the book, the general situation can be described by an outcome variable Y , error-free covariate(s) Z and unobservable covariate(s) X , for which proxies W are available. In the simplest setting $W=X+\text{error}$. Interest focuses on the (non-linear) regression of Y on X and Z . If X is replaced by its proxy W , a biased estimate is obtained of the regression coefficient of X . In the simple situation of a one-dimensional X and no Z , the regression coefficient is shrunken towards zero and it has to be corrected for this attenuation. This phenomenon is has long been known in psychometrics, where error-free scales are very rare.

Throughout the book it is assumed that either the (co)variance of the measurement error in W is known or that an instrumental variable is present, that is a second proxy for X with error independent of the error in W . To my surprise the well-known orthogonal regression is discarded early in the book. Orthogonal regression is based on the assumption that the ratio of the errors in Y and X as known (usually, it is taken to be equal to one). That makes sense in the symmetric situation where Y and X are stemming from different instruments measuring the same quantity. However, in pure regression modeling there is no such symmetry. In the regression model, there is not only measurement error for Y , but also the 'equation error' arising from lack of fit of the model. Since the latter error is usually of a different order than the measurement error, orthogonal regression leads to more attenuation than necessary.

The book contains fourteen chapter and an appendix with a review of the statistical theory needed throughout the book. It has a logical set up from simple to complicated

techniques for estimating the regression equations. The last three chapters are a bit astray. A topic as density estimation and deconvolution, discussed in chapter twelve, may be mathematically related to the main theme of the book, but it is of limited importance for the practitioner dealing with measurement error.

The first two chapters are purely introductory. They introduce the notation mentioned above and explain the problem on the example of simple linear regression. The next chapters discuss different solutions. In chapter three regression calibration is introduced. The idea is first to estimate X from W and then to use the estimated X in fitting the regression equation. Practically, this gives reasonable results, although the method is only approximately consistent. Chapter four deals with the SIMEX methodology, standing for SIMulation- EXtrapolation. This is a very recent idea due to Stefanski and Cook and further developed by Carroll and others. The elegant idea is to add extra noise to W by simulation. A plot is made of the regression coefficients versus the amount of added noise. By extrapolating this plot to the left, an estimate can be obtained for the noiseless situation where noise is 'subtracted' instead of added. The idea is charming and well illustrated by examples. Chapter five discusses the use of instrumental variable. This is closely related to the regression calibration of chapter three. In chapter six more advanced methods are introduced based on the construction of unbiased estimating equations yielding consistent estimates. The methods are practically feasible, but are not very well explained. I found this part hard to read. In chapters seven and eight the switch is made from functional to structural models. While functional models treat the X 's as completely unknown, the structural approach is based on the assumption of an underlying distribution for X . Under that assumption classical statistical methods can be used: likelihood methods in chapter 7 and Bayesian methods including the Gibbs Sampler in chapter eight. Chapter nine discusses semi-parametric methods for the situation that the true X is observed for part of the sample. This part is quite theoretical and no examples are given. The application oriented reader will lose track here. The remainder is mainly theoretic and I will not discuss it in detail.

This book contains interesting material that is relevant to all applied statisticians having to deal with measurement error in explanatory variables. The authors operate at the frontier of statistical science. Having appeared in 1995 it already refers to papers published in 1996. The book contains many examples and it also gives the electronic mailbox where the necessary software can be obtained.

To summarize my review, I think this is an important book of high theoretical level and great practical potential. However, it is no cook book. The practitioner, who thinks that he will get easy recipes about how to handle his own practical problems, might be disappointed when he lays hands on this book. It is no easy reading and it requires a thorough study to find out what the authors exactly mean. I think that is inevitable for a book that contains so many new ideas and methods. The newest material is there and the applied statistician is challenged to find in the book what he needs. That is no

sinecure, but it is very worth trying.

prof. dr H.C. van Houwelingen
Vakgroep Medische Statistiek
Rijksuniversiteit Leiden

Davidian, M., Giltinan, D.M.

Nonlinear models for repeated measurement data

1995, Monographs on statistics and applied probability 62,
Chapman & Hall, London,
359 blz, £ 32.00, ISBN 0-412-98341-9.

Dit is een leerzaam en goed geschreven leerboek over, de titel zegt het al, niet-lineaire modellen voor herhaalde metingen. Bij herhaalde metingen wordt elk onderzocht individu (proefpersoon, proefdier, of wat voor eenheid dan ook) herhaaldelijk gemeten. In de klassieke situatie is dat op vaste, voor elk individu gelijke meetmomenten, en wordt de afhankelijkheid van verklarende variabelen door een lineair model uitgedrukt. Daarvoor worden gewoonlijk traditionele multivariate of mixed-model methoden gebruikt.

Dit boek gaat echter over designs waarin de meetmomenten voor de individuen verschillend kunnen zijn, en over modellen waarbij zowel de binnen-individuele als de tussen-individuele afhankelijkheid van verklarende variabelen door niet-lineaire functies worden uitgedrukt. Voor de kansverdelingen van de uitkomsten wordt de aandacht gericht op continue (bijv. multivariaat normale) verdelingen. Discrete afhankelijke variabelen worden niet behandeld.

Het bestudeerde *General Hierarchical Nonlinear Model* is als volgt gedefiniëerd: voor individu i wordt een $n_i \times 1$ vector van uitkomsten

$$\mathbf{y}_i = \mathbf{f}_i(\boldsymbol{\beta}_i) + \mathbf{e}_i \quad (1)$$

waargenomen. Hierbij staat $\mathbf{f}_i(\boldsymbol{\beta}_i)$ voor een vector met componenten $f(x_{ij}, \boldsymbol{\beta}_i)$, en heeft de vector van residuen $\mathbf{e}_i$ verwachting 0 en covariantiematrix $\mathbf{R}_i(\boldsymbol{\beta}_i, \boldsymbol{\xi})$. Voor de individu-afhankelijke coëfficiënten $\boldsymbol{\beta}_i$ geldt het model

$$\boldsymbol{\beta}_i = \mathbf{d}(\mathbf{a}_i, \boldsymbol{\beta}, \mathbf{b}_i), \quad (2)$$

waarbij de $\mathbf{b}_i$ onafhankelijke en gelijk verdeelde vectoren zijn. Hierin zijn verder de x_{ij} en de $\mathbf{a}_i$ verklarende variabelen. De statistische parameters zijn $\boldsymbol{\beta}$ en $\boldsymbol{\xi}$.

Als er bijvoorbeeld een onderliggende tijdsdimensie is, dan kunnen de x_{ij} functies van de tijd zijn (i.h.a. betreffen deze variabelen de j 'e meting voor individu i), terwijl de $\mathbf{a}_i$ constante covariaten zijn. Ook aan variërende covariaten wordt aandacht besteed. Als de functies f en $\mathbf{d}$ lineair worden verondersteld, en als voor de residuen $\mathbf{e}_i$ en de stochastische coëfficiënten wordt verondersteld dat ze normaal verdeeld zijn, wordt het hiërarchische lineaire model verkregen. Dit is een speciaal geval van het gemengde lineaire model, en wordt gekenmerkt door geneste stochastische coëfficiënten.

Het boek is als een prettig en up-to-date leerboek opgezet. Het begint met motiverende voorbeelden, een beantwoording van de vraag wat je zou doen als je data van één individu had (niet-lineaire regressie), en een overzicht van de bekende statistische methoden voor het hiërarchische lineaire model. De hoofdmoot van het boek bestaat uit vier hoofdstukken die elk een andere benadering behandelen voor het schatten en toetsen van de parameters in het hierboven aangegeven model. Die benaderingen zijn, achtereenvolgens:

1. tweetraps-procedure: inferentie op grond van schatting van de individuele coëfficiënten β_i in (1), waarna deze worden geanalyseerd volgens (2), rekening houdend met de onnauwkeurigheid van de schattingen;
2. linearisatie (Taylor-reeksen) van de functies f en d ;
3. non-parametrische en semi-parametrische inferentie wat betreft de verdeling van de b_i ;
4. Bayesiaanse inferentie met behulp van Gibbs sampling.

Deze vier hoofdstukken zijn afzonderlijk leesbaar, en zijn voor praktische doeleinden behoorlijk ver uitgewerkt (inclusief opmerkingen over modelselectie, rekenpartijen, software, e.d.). Het boek wordt afgesloten met drie hoofdstukken met uitgebreide toepassingen (de eerste twee hiervan gaan over biostatistische, het laatste gaat over andere toepassingen).

De auteurs richten zich, blijkens het voorwoord, vooral op biostatistici. Maar ook statistici in andere gebieden kunnen veel van dit boek leren. De nadruk op biostatistiek kan enerzijds worden gezien tegen het licht van het feit dat de auteurs zelf in dit gebied werkzaam zijn en dat het boek een overzicht geeft van veel recent werk (van de auteurs en van anderen) dat vooral in biostatistische tijdschriften verschenen is. Anderzijds gaat het model er van uit dat de functies f en d bekend zijn. Dat trekt een wissel op theoretische kennis die in biologische toepassingen misschien vaker gedekt is dan in andere toepassingsgebieden. Verder geven de auteurs, begrijpelijkerwijs, meer aandacht aan bijdragen vanuit de biostatistische dan aan die uit andere hoek (bijv. wordt niet genoemd het werk van H. Goldstein en medewerkers over tweede-orde Taylorreeksen benaderingen, noch het gemak dat men bij multivariate afhankelijkten kan hebben van modellen met drie hiërarchische niveaus).

Een leuk aspect van het boek zijn de discussieparagrafen die in de meeste hoofdstukken staan, en waarin vergelijkingen worden gemaakt met wat elders in de literatuur staat, keuzes van de auteurs worden beargumenteerd, voor- en nadelen van benaderingswijzen worden besproken, enz.

Ik kan het boek aanraden aan ieder die belangstelling heeft voor niet-lineaire herhaalde metingen en zich in staat acht om door de formules de ideeën te blijven zien.

*prof. dr T.A.B. Snijders
Vakgroep Statistiek en Meettheorie
Rijksuniversiteit Groningen*

Waterman, M.S.

Introduction to computational biology

1995, Interdisciplinary statistics,

Chapman & Hall, London,

431 blz, £ 29.99, ISBN 0-412-99391-0.

Hoewel 'computational biology' doet vermoeden dat het gaat om alles in de biologie waar rekenen aan te pas komt, staan deze twee woorden meestal voor een veel gespecialiseerder gebied. Met 'computational biology' wordt namelijk de combinatie van wiskunde en biologie aangeduid die zich bezighoudt met moleculair biologische data, in het bijzonder die van DNA- en RNA-sequenties, chromosomen, eiwitten en dergelijke. Dit onderwerp staat momenteel bij statistici zeer in de belangstelling. Immers, door de enorme ontwikkelingen in de moleculaire biologie van de afgelopen decennia en de gigantische hoeveelheden data die ten gevolge hiervan worden gegenereerd, zijn er voor wiskundigen en met name statistici vele interessante en uitdagende problemen op dit gebied. Het boek van Waterman is het eerste uitgebreide boek over computational biology.

Het boek behandelt vooral wiskundige aspecten van DNA- en RNA-sequenties en van chromosomen, maar laat eiwitten grotendeels buiten beschouwing. De auteur beoogt de "fascinating structure of biological data and problems" te introduceren bij mensen met een wiskundige achtergrond. Mijns inziens slaagt hij daar zeker in: het boek is met veel enthousiasme en kennis van zowel biologische als wiskundige zaken geschreven.

Aangezien het om toepassing van de wiskunde op moleculair biologische problemen gaat, is enige basiskennis van de moleculaire biologie noodzakelijk. In het boek is een korte introductie hierin opgenomen, maar die is aan de summiere kant. Daar de wiskunde in het boek gemotiveerd wordt door de biologische problemen, worden er namelijk erg veel vaktermen gebruikt die niet in de introductie aan de orde komen. Dit is, bijvoorbeeld, het geval wanneer om tot een goed wiskundig model te komen, uitgelegd wordt hoe bepaalde data verkregen zijn. De beschrijving van het biologische experiment, hoewel helder, zal de gemiddelde wiskundige boven de pet gaan. Kennis van de terminologie is echter niet noodzakelijk om de wiskunde te kunnen volgen, maar maakt het geheel natuurlijk wel interessanter. Voor wie bereid is over onbekende biologische details heen te lezen, slaagt de auteur er toch wel in om van elk probleem de biologische achtergrond enigszins begrijpelijk te schetsen, duidelijk het biologische probleem en het daarvan afgeleide wiskundige probleem te formuleren.

De lezer wordt geacht basiskennis te hebben van kansrekening, statistiek en calculus, wat me correct lijkt. De overige wiskunde die aan de orde komt, zoals, bijvoorbeeld, grafentheorie of simulated annealing, wordt uitgebreid geïntroduceerd. Enige kennis van informatica—de auteur noemt de begrippen algoritme en complexiteit—is ook noodzakelijk, evenals ervaring met het gebruik van computers.

Een uitgebreid scala aan biologische problemen komt aan de orde: het maken en bestuderen van restrictiekaarten van DNA, de representativiteit van clone libraries, het

construeren van genenkaarten, het lezen van DNA-sequenties, verschillende aspecten van het vergelijken van sequenties (met name het traceren van overeenkomstige patronen in twee of meer sequenties), het voorspellen van de secundaire structuur van RNA gegeven de lineaire sequentie en tenslotte het afleiden van evolutionaire relaties tussen groepen organismen uit hun DNA- of eiwitsequenties. Tot slot is in het boek nog een handig hoofdstuk opgenomen met achtergrondliteratuur, zowel over de biologie als over de wiskunde, alsook enkele verwijzingen naar verwant lopend onderzoek. Elk hoofdstuk besluit met enkele opgaven.

Bovenstaande, evenals trouwens de inhoudsopgave van het boek, wekt misschien de indruk dat het boek meer geschikt is voor (theoretisch) biologen dan voor wiskundigen. Dat is echter zeker niet het geval. Na een korte uitleg van het biologische probleem, wordt het in het algemeen al snel zeer wiskundig. De gebruikte stochastische modellen worden geïntroduceerd, evenals de wiskundige technieken—vaak combinatorische, vaak statistische—en de computeralgoritmen. Doordat de wiskunde gemotiveerd wordt door de biologie, is de gepresenteerde wiskunde zeer veelzijdig en zijn de behandelde wiskundige problemen van een boeiende complexiteit. Kortom, een aanrader voor elke wiskundige die in bovengenoemde biologische problemen is geïnteresseerd!

*dr M.C.M. de Gunst
Vakgroep Wiskunde
Vrije Universiteit Amsterdam*

Guttorp, P.

Stochastic modelling of scientific data

1995, Stochastic modelling series,
Chapman & Hall, London,
372 blz, £ 29.95, ISBN 0-412-99281-7.

De nadruk in dit boek ligt op het modelleren en analyseren van data met behulp van methoden uit de kansrekening en statistiek. De lezer dient te beschikken over basis-kennis uit deze vakgebieden alsmede over enige kennis van analyse en lineaire algebra. Kennis van stochastische processen, i.h.b. van Markovketens is gewenst, maar strikt genomen niet noodzakelijk.

Uitgangspunt van het boek is zoals gezegd het modelleren en analyseren van data. Hierbij wordt relatief veel moeite gedaan om de fysische of biologische achtergronden van allerlei fenomenen die in een groot aantal voorbeelden aan de orde komen uit te leggen om zo de lezer vertrouwd te maken met deze uit de (statistische) praktijk afkomstige problemen. Veel minder diepgaand wordt ingegaan op de wiskundige onderbouwing van de probabilistische eigenschappen van een gepresenteerd model en op bijv. eigenschappen van geponeerde schatters.

In grote trekken zien we dat elk hoofdstuk is opgebouwd volgens een min of meer vast stramien, eerst wat algemene wiskundige theorie, dan wat toepassingen en voorbeelden, die vaak verder in hetzelfde hoofdstuk (of in een volgend) terugkeren, wanneer de

theorie wat verder uitgediept wordt. Dit werkt motiverend voor iemand, die wil weten waarom je bepaalde technieken zou kunnen gebruiken in een praktische context. Dit is duidelijk een boek voor mensen die in de praktijk werkzaam (willen) zijn.

Laten we even de inhoud van het boek nader bekijken en daarbij de hoofdstukindeling aanhouden. Hoofdstuk 1 is zeer inleidend van karakter en geeft wat filosofische getinte overwegingen voor het gebruik van stochastische modellen, bijv. om te beschrijven dan wel te voorspellen.

Hoofdstuk 2 gaat over Markov processen in discrete tijd. Deze worden behandeld aan de hand van een groot aantal voorbeelden (neerslagcijfers in Snowquallmie Falls, stralingsschade, Ehrenfest diffusiemodel, Hardy-Weinberg model in genetica, een model voor de volgorde klinker/medeklinker in Evgenij Onegin, verspreiding van een epidemie, een neurofysiologisch model voor ion-kanalen). Eigenschappen van Markovketens (stationair, ergodisch gedrag, classificatie van toestanden etc) worden aan de hand van deze voorbeelden geïntroduceerd. Tevens vinden we in dit hoofdstuk statistische analyse van parameterschatters voor Markovketens, Likelihood theorie en Monte Carlo methoden.

In hoofdstuk 3 vinden we Markov processen in continue tijd, Kolmogorov vergelijkingen, statistische inferentie voor deze modellen en wederom veel uitgebreide voorbeelden waarvan we noemen: geboorte/sterfte en immigratie/emigratie gedrag in een bavianenpopulatie in Kenya, het modelleren van neurale activiteit, vorming van bloedcellen.

Hoofdstuk 4 is gewijd aan stochastische velden, welke geïntroduceerd worden naar aanleiding van het Ising model voor ferromagnetisme, modellen voor de verspreiding van een ziekte onder planten, een genetisch model voor een Eskimo populatie, een model voor A/B/O classificatie voor bloedgroepen.

Hoofdstuk 5 gaat over puntprocessen. Het eerste gepresenteerde voorbeeld gaat over het tellen van auto's die in clusters een zeker punt langs een snelweg passeren. Een ander voorbeeld is het modelleren van voltagepieken in een neuron. Op het theoretische vlak vinden we een behandeling van begrippen als intensiteit, likelihood voor puntprocessen, het schatten van (parameters in) de intensiteit, het vinden van een model voor de intensiteit (zoals het Cox-model). Verder zien we ook ruimtelijke puntprocessen, zoals in een model voor de locatie van bomen in een tropisch regenwoud.

Tenslotte is het onderwerp van Hoofdstuk 6 Brownse beweging en diffusie. Elementaire eigenschappen van de Brownse beweging (verdeling van de aangroeiingen, begrensde kwadratische variatie etc.) komen aan de orde, de Fokker-Planck vergelijking wordt afgeleid. Verder zien we een elementaire uiteenzetting van stochastische integralen, de Ito formule wordt geponeerd, het Ornstein-Uhlenbeck proces en geometrische Brownse beweging. Er wordt iets gezegd over de likelihood schatter in de context van stochastische differentiaalvergelijkingen. Toepassingen die gegeven worden bevatten een stochastisch model voor het aantal cellen van een specifiek type dat informatie bevat over de kwaliteit van het immuunsysteem. Dit is van belang voor het analyseren van de ontwikkeling van AIDS.

Elk hoofdstuk wordt besloten met een collectie vraagstukken (waar en passant nog wel eens wat nieuwe methodologie in verstopt zit, zoals het Kalman filter of de boots-trap) die in drie categorieën uiteenvalt, de theoretische vraagstukken, de "computational excercises", waarbij de hulp van een computer vaak onmisbaar is en vraagstukken waarbij een verzameling data geanalyseerd dient te worden.

Afsluitend kunnen we zeggen dat dit een boek is waarin de modellen centraal staan en waarin allerlei wiskundige technieken die bij de analyse hiervan van pas komen als technisch gereedschap gepresenteerd worden. Het boek is hierdoor geen boek over een wiskundige theorie, maar wel een waaruit je kunt leren hoe je met gegevens uit de praktijk op een wiskundig verantwoorde wijze kunt omgaan. Het grote aantal voorbeelden, die uitgebreid behandeld worden, is een van de sterke kanten van dit boek.

dr P.J.C. Spreij

Vakgroep Econometrie

Vrije Universiteit Amsterdam

Shiryaev, A.N.

Probability

1995, Graduate texts in mathematics,

Springer-Verlag, Berlin,

xiv + 609 blz, DM 89.00, ISBN 0-387-94549-0.

In dit boek vinden we zo ongeveer alles wat van fundamenteel belang is in de theorie van stochastische processen met een discrete tijdsparemeter, zoals wetten van de grote aantallen, zwakke convergentie, Markov ketens, martingalen, zwak stationaire processen, om maar een greep uit de inhoud te doen. Voor een preciese behandeling van deze onderwerpen is het gebruik van maattheorie onmisbaar. Dit is grotendeels de inhoud van hoofdstuk 2, waarin we tevens typische toepassingen vinden in de kansrekening zoals Kolmogorovs consistentiestelling.

Echter zolang de stochastische variabelen maar eindig veel waarden aannemen valt er van alles te doen zonder expliciet gebruik van maattheorie, omdat de benodigde σ -algebra's in de regel door een eindige partitie gegenereerd worden. In zo'n geval kun je bijvoorbeeld de voorwaardelijke verwachting definiëren zonder een beroep op de stelling van Radon-Nikodym te doen. Een uitgebreid eerste hoofdstuk (Elementary Probability Theory) behandelt een groot deel van de hierboven genoemde onderwerpen langs deze lijn. Lang niet alles is trouwens elementair.

Hoofdstuk 3 is gewijd aan convergentie van kansmaten. We zien hier klassieke onderwerpen als de stelling van Prohorov, metrieken voor zwakke convergentie, maar ook een uitstapje naar het gebruik van Hellingerintegralen om bijvoorbeeld contiguiteit te karakteriseren, een tamelijk recente onderzoeksinteresse van de auteur. Ik vroeg me af of dit een natuurlijke plaats heeft in dit boek gezien de keuze van de overige onderwerpen.

De korte hoofdstukken 4 en 5 behandelen respectievelijk *iid* rijtjes en basisbegrippen uit de ergodentheorie. Hoofdstuk 6 over zwak-stationaire processen is weer wat

uitgebreider met onder meer spectraaltheorie en (een generalisatie van) het Kalmanfilter.

In hoofdstuk 7 komen martingalen aan de orde, inclusief de bekende convergentiestellingen, maar ook onderwerpen als de Kakutani dichotomie (slordig gezegd zijn twee oneindige producten van onderling equivalente kansmaten equivalent of orthogonaal) en een Itô-formule in discrete tijd, wederom een hobby van de auteur. Het laatste hoofdstuk gaat over Markovketens, classificatie van toestanden en stationaire verdelingen.

Zoals uit deze opsomming blijkt hebben we te maken met een rijk boek. De behandelde onderwerpen zijn voor het grootste gedeelte klassiek, maar nu en dan is de persoonlijke smaak van de auteur te herkennen. Verder vereist het begrijpen van de tekst in het algemeen geen specifieke voorkennis anders dan standaardbegrippen uit de analyse met als uitzondering de spectraaltheorie voor zwak-stationaire processen, waar wat theorie van Hardyruimten bekend wordt verondersteld. Hierdoor en door de verscheidenheid aan onderwerpen zou het boek zeer geschikt zijn om een college voortgezette kansrekening uit te geven. Een waarschuwing is echter op zijn plaats. Er staan, zeker voor een tweede druk, naar mijn smaak veel te veel slordigheden in. Drukfouten, inconsistente notatie, hier en daar een rommelige bewijsvoering en foutieve verwijzingen zijn nogal eens hinderlijk bij de bestudering. Ook leidt het streven naar af en toe zeer algemene resultaten niet altijd tot een transparante presentatie.

*dr P.J.C. Spreij
Vakgroep Econometrie
Vrije Universiteit Amsterdam*

Rossmann, A.J.

Workshop Statistics

1996, Springer-Verlag, Berlin,
452 blz, DM 44.00, ISBN 0-387-94497-4.

Het boek geeft een introductie op het gebied van statistiek. Doelgroep hierbij zijn studenten in het hoger onderwijs. De auteur legt veel nadruk op actief leren en het begrijpen van concepten. Studenten maken veel gebruik van 'echte' data met betrekking op allerlei gebieden en werken met rekenmachines/computers.

Inhoud: data, variabelen, maten van ligging en spreiding, beschrijving en uitbeelding van verbanden, normale verdeling, correlatie, regressie, betrouwbaarheidsinterval, significantietoetsen.

Opbouw van een hoofdstuk:

- overview: het onderwerp wordt kort geplaatst in de context;
- objectives: doelen die student heeft bereikt na het succesvol doorlopen van het hoofdstuk;

- preliminaries: aantal vragen die studenten vóór de les beantwoorden; de vragen hebben als doel de student alvast te laten nadenken over een aantal zaken, soms wordt gevraagd om wat gegevens te verzamelen;
- in-class activities: opdrachten die studenten (in groepjes) tijdens de les uitvoeren
- homework activities: opdrachten die de studenten thuis maken ter controle en bezinking;
- wrap-up: korte samenvatting.

De opzet past goed bij het idee dat studenten pas echt iets leren als ze zelf actief aan de slag gaan. De docent heeft bij het gebruik van dit boek andere didaktische methoden nodig, iets waar Rossman heel kort iets over zegt. De naamgeving aan 'activities' hangt volgens mij meer samen met de wens die de auteur heeft over de werkhouding van de studenten, dan met de inhoud van het omschrevene. Vaak is een 'activiteit' gewoon een 'opdracht'. De onderwerpen van de activiteiten zijn vaak amerikaans. Dat zal in sommige gevallen de studenten minder aanspreken, en dan is het voordeel van het gebruik van 'echte' data nihil. Het 'vertalen' naar de nederlandse situatie kan een (tijdsintensieve) oplossing zijn.

Een voorbeeld van (een deel van) de gang van zaken in een hoofdstuk (onderwerp regressie):

- preliminaries:
 - Als één stad verder is dan een andere, verwacht je dan dat het over het algemeen méér kost om naar de verdere stad te vliegen?
 - Doe een gok hoeveel (gemiddeld) elke extra kilometer toevoegt aan de totale prijs van het vliegticket.
- in-class activities:
 - (oorspronkelijk gegeven een tabel met de kosten van vliegtickets van Baltimore naar diverse steden in de USA, ik heb dat vertaald - met dank aan een reisbureau - in kosten vanaf Amsterdam naar diverse steden in Europa) Maak een spreidingsdiagram van de gegevens en trek die rechte lijn die volgens jou het beste het verband weergeeft.
 - Bepaal de vergelijking van die lijn. Wat is de richtingscoëfficiënt van die lijn? Wat kost een extra kilometer meer? Vergelijk met preliminaries.
 - Bepaal de kleinste kwadratenlijn op basis van formules en de waarnemingen. Teken deze lijn in het spreidingsdiagram. Wat is de richtingscoëfficiënt van die lijn? Wat kost een extra kilometer meer? Vergelijk met preliminaries.

Studenten die nog niet eerder zelfstandig met statistiek leren zijn bezig geweest, zullen in het begin wat vreemd opkijken. Commentaar dat ik kreeg van studenten was bijvoorbeeld n.a.v. de tweede preliminary: 'ik weet toch niet wat dat kost, hoe kan ik daar dan antwoord op geven?'. Het is erg nieuw voor ze dat gevraagd wordt naar hun mening, in plaats van naar 'het goede antwoord'. Natuurlijk zijn er ook veel opdrachten waar een precies antwoord (getal) verwacht wordt. De antwoorden van die vragen zijn niet opgenomen. Waarschijnlijk is dat bedoeld om het onderling overleg te bevorderen. Na het doorlopen van de opdrachten valt vooral de samenhang in de opdrachten op. Vaak wordt terugverwezen naar een eerdere activiteit. Studenten kunnen zo kijken of ze in het begin een beetje in de goede richting zaten met hun gok, de samenhang tussen de onderwerpen wordt hierdoor ook vergroot.

Na het bekijken van het boek bleef nog één vraag over: Waarom zijn de pagina's uitscheurbaar en voorzien van drie gaten?

*ir.M.M. Jansen
sector Bedrijfskader
Hogeschool Eindhoven*

Valkenburg, H.A., Hofman, A.

Een kwart eeuw Hippocratische epidemiologie

1995, Bunge, Utrecht,
95 blz, fl 27,50, ISBN 90-6348-478-X.

De beide auteurs zijn emeritus-hoogleraar en hoogleraar Epidemiologie te Rotterdam. Het boek begint met de afscheidsrede van Dr. H.A. Valkenburg (1989) en de oratie van Dr A. Hofman (1992). Daarna volgen beide heren met een persoonlijke geschiedschrijving van de epidemiologie in Rotterdam.

Dr. H.A. Valkenburg start met de beschrijving van de jaren 1970-1988, de tijd dat hij hoogleraar was in Rotterdam. Het is een persoonlijk verhaal dat ook ingaat op achtergronden achter de keuze voor een epidemiologische leerstoel in Rotterdam. Ook de opbouw van het epidemiologisch instituut komt uitgebreid aan de orde. Bijzonder aandachtspunt hierbij is het Epidemiologisch Preventief Onderzoek Zoetermeer (EPOZ), een langlopend onderzoek waarbij ruim 10 000 respondenten betrokken waren.

H. Hofman bespreekt de periode 1977-1995. Zijn verhaal beperkt zich meer tot de onderwerpen van onderzoek in dit vakgebied.

Het boek eindigt met een beschrijving van de onderzoeken die in deze vijftientig jaar gedaan zijn in Rotterdam. Zij zijn hiervoor ingedeeld in de volgende 10 thema's:

1. Artherosclerose
2. Bewegingsapparaat
3. Epidemiologische tekstboeken
4. Genetische epidemiologie

5. Hart- en vaatziekten
6. Hypertensie
7. Infectieziekten
8. Kanker
9. Neurologische ziekten
10. Pediatrische epidemiologie

Tot slot worden alle promovendi genoemd die gepromoveerd zijn in de periode 1970- (medio) 1995 en een promotor hadden uit de epidemiologische groep. Hierbij wordt van de 46 promovendi de promotiedatum, de titel van het proefschrift en de betrokken promotor genoemd. Ook wordt van elk van hen de huidige functie vermeld.

De doelgroep van het boek is niet direct duidelijk. Het is geen studieboek of wetenschappelijke verhandeling, en kan ook niet als zodanig gelezen worden. Voor onderzoekers op dit gebied is het meeste bekend. A. Querido, bouwdecaan van de medische faculteit in Rotterdam karakteriseert het boek als een visitekaartje van de groep, aangeboden aan de overheid.

Het boek is prettig leesbaar, en geeft een duidelijk overzicht van het vakgebied zoals het aan de Erasmus universiteit wordt beoefend. Ook (en misschien wel vooral) geschikt voor niet-epidemiologen.

*E. Jeurissen
Obistat*

Longford, N.T.

Models for uncertainty in educational testing

1995, Springer series in statistics,
Springer-Verlag, Berlin,
300 blz, DM 68,00, ISBN 0-387-94513-X.

Educational Testing Services (ETS) is een van de grote testbureau's in de Verenigde Staten. Daar wordt hetzelfde soort werk gedaan als op het Nederlandse CITO. Jaarlijks worden allerlei soorten toetsen voor gebruik in het onderwijs gemaakt en afgenomen, om de prestaties van leerlingen en studenten te meten. Veel grootschalig onderwijskundig onderzoek in de Verenigde Staten, zoals het National Assessment of Educational Progress (NAEP) programma wordt door het ETS begeleid. Het zal duidelijk zijn dat de constructie van toetsen, die in het onderwijs gebruikt worden en de analyse van de gegevens een grote diversiteit aan statistische en psychometrische problemen met zich meebrengt. De auteur van dit boek werkt op het ETS en beschrijft in dit boek een groot aantal van dergelijke problemen. In de psychometrie zijn er globaal drie verschillende benaderingen om de gegevens van prestatietoetsen te analyseren. Dat zijn

het klassieke testmodel, dat uitgaat van een ware score definitie, de generaliseerbaarheidstheorie, dat uitgaat van variantiecomponenten en itemresponstheorie (IRT), dat de kans op een respons modelleert aan de hand van itemparameters. Hoewel de laatste twee decennia de itemresponstheorie zich mag verheugen in een enorme belangstelling van psychometrici en onderzoekers in de onderwijskunde, wordt in dit boek de benadering van de testgegevens met behulp van de klassieke testtheorie en de generaliseerbaarheidstheorie ondersteunt. Dit is wellicht jammer, omdat itemresponstheorie geenszins meer in de kinderstoelfase verkeert en diverse ontwikkelingen praktisch toepasbaar zijn geworden. Dit blijkt wel uit de recente uitgaven van twee (hand) boeken over IRT modellen van Fischer en Molenaar (1995) en Van der Linden en Hambleton (1996).

Toch vind ik dit boek een aanwinst. Door het grote aantal voorbeelden uit de onderwijskundige praktijk geeft Longford aan dat veel problemen met testgegevens op ETS opgelost worden zonder IRT modellen. Hoewel hij geen echte vergelijking tussen beide benaderingen geeft, laten de voorbeelden voor mij zien dat ook uitstekende resultaten verkregen kunnen worden zonder de IRT benadering.

Dit boek heeft een evenwichtige opbouw. Begrippen worden duidelijk en evenwichtig, ondersteund met formules, uitgelegd. Waar nodig worden statistische begrippen, zoals de jackknife, shrinkage schatters, restricted ML, en small area predictie uitgelegd. In elke hoofdstuk worden een aantal voorbeelden gepresenteerd om de technieken te illustreren. Veel onderwijskundige onderzoekers en sociale wetenschappers die testgegevens willen analyseren, kunnen dit boek raadplegen voor de oplossing van hun problemen. Ook de wat meer statistisch en psychometrisch geschoolde onderzoekers zullen dit boek als naslagwerk kunnen appreciëren.

De hoofdstukken 2 tot en met 4 behandelen de analyse van subjectieve beoordelingen van expert-beoordelaars. In hoofdstuk 2 ligt het accent op de schatting van de betrouwbaarheid van beoordelingen via variantie-componenten. In hoofdstuk 3 wordt met behulp van shrinkage schatters aangegeven hoe de ware scores van respondenten beter geschat kunnen worden en in hoofdstuk 4 wordt nader ingegaan op het analyseren van beoordelingen die verkregen zijn van verschillende essays. In hoofdstuk 5 wordt ingegaan op adequate samenvatting van itemkenmerken, zoals itemmoeilijkheid en het verschijnsel dat items verschillend worden beantwoord door subpopulaties (DIF). Wanneer verschillende tests verondersteld worden dezelfde trek te meten, dan kunnen de scores op deze tests op eenzelfde schaal geplaatst worden. Dit wordt equivalenten genoemd, en hoofdstuk 6 behandelt dit. In hoofdstuk 7 wordt ingegaan op het analyseren van testgegevens uit grootschalige onderwijskundig onderzoek, zoals het NAEP programma in de VS. Enkele problemen rond statistische inferentie en steekproeftrekkingen worden hier behandeld. Statistische inferentie voor small-areas aan de hand van de gegevens uit grootschalige surveyonderzoek, inclusief small-area predictie wordt in hoofdstuk 8 besproken. Tenslotte wordt via logistische regressie het probleem van het adequaat opstellen van zak/slaag grenzen uitgelegd en in hoofdstuk 10 wordt de analyse van onvolledige longitudinale data voor de groei van kennis aan de hand van het onvermijdelijke EM-algoritme besproken.

Als conclusie wil ik stellen dat dit boek een aanvulling is voor de vele psychometrische boeken die de analyse van testgegevens vanuit de IRT benaderen.

Verwijzingen:

- [1] Van der Linden, W.J. en Hambleton, R.K. (1996). Handbook of modern item response theory. New York: Springer Verlag.
- [2] Fischer, G.H. en Molenaar, I.W. (1995). Rasch models: foundations, recent developments and applications. New York: Springer Verlag.

*prof. dr M.P.F. Berger
Vakgroep Methodologie en Statistiek
Rijksuniversiteit Limburg*

Roes, C.B.

Shewhart-type charts in statistical process control

1995, Thesis Publishers, Amsterdam,
121 blz, fl 35,-, ISBN 90-5170-365-1.

Dit boek bevat het proefschrift van de auteur. Het behandelt verscheidene praktische aspecten van het gebruik van Shewhart regelkaarten binnen de statistische kwaliteitszorg. De auteur behandelt problemen bij het toepassen van Shewhart regelkaarten die hij tegenkwam in zijn industriële praktijk.

Hoofdstuk 1 is bedoeld als inleiding en geeft een overzicht van recent onderzoek in Nederland op het gebied van statistische kwaliteitszorg.

De theoretische fundering van het boek is te vinden in hoofdstuk 2, waar een zeer algemeen kader gegeven wordt voor Shewhart regelkaarten, gebaseerd op [2].

De overige hoofdstukken zijn gewijd aan praktische problemen die optreden bij het gebruik van Shewhart regelkaarten. Hoofdstuk 3 behandelt Shewhart regelkaarten voor individuele waarnemingen. Het probleem hierbij is het schatten van varianties. Verschillende methoden worden met elkaar vergeleken en beoordeeld. Hoofdstuk 4 is gebaseerd op [4] en bevat een ander praktisch probleem, nl. proces-inherente variantie. Dit wordt aangepakt met behulp van variantie-analyse met zowel vaste als stochastische effecten. Tenslotte komen in hoofdstuk 5 problemen die optreden bij laag-volumeproductie aan de orde. Dit hoofdstuk is gebaseerd op [1] (verg. [3]).

Het boek is prettig leesbaar geschreven, al rammelt het Engels wel op een aantal plekken. Stellingen en bewijzen worden apart behandeld in appendices. Elk hoofdstuk wordt besloten met concrete adviezen voor de praktijk. Naast theoretische gedeeltes zijn er ook verhelderende praktijkgevallen. Gezien de vereiste voorkennis lijkt het boek bestemd voor mathematisch statistici en ervaren bedrijfsstatistici. Deze twee groepen zullen veel interessante zaken aantreffen in dit boek. Het boek vult een leemte in de literatuur die voornamelijk bestaat uit op praktijk gerichte boeken zonder enige (theoretische) diepgang en theoretische boeken zonder oog voor de praktijk.

Tenslotte wil ik nog enige minpunten noemen van dit boek noemen. Persoonlijk ben ik niet zo gecharmeerd van proefschriften die bestaan uit losse artikelen. Dit leidt

tot onnodige overlap en het ontbreken van een uniforme lijst met verwijzingen aan het eind. Hoewel de auteur zich wel enige moeite heeft getroost, was het in dit tijdperk van tekstverwerkers een kleine moeite geweest voor de auteur om dit te vermijden en tevens een index toe te voegen. Een statistisch punt van kritiek is de aanname van normaliteit. Het is vreemd dat in een boek waarin zoveel aandacht wordt besteed aan modelaannamen, slechts één pagina (pp. 15-16) aan normaliteit wordt besteed. Er is veel meer bekend over bijv. verdelingsvrije regelkaarten dan de auteur doet vermoeden.

Samenvattend: een interessant boek voor geïnteresseerden in de statistische kwaliteitszorg (mits men voldoende bagage bezit).

Verwijzingen:

- [1] R.J.M.M. Does, K.C.B. Roes en A. Trip, A general set-up for designing control charts in a mixed model, Mathematical preprint series 95-09, Universiteit van Amsterdam, 1995.
- [2] R.J.M.M. Does en B.F. Schriever, Variables control chart limits and tests for special causes, Stat. Neerl. 46 (1992), 229-245.
- [3] K.C.B. Roes, SPC bij laag volume en kleine series fabricage, Kwant. Meth. 53 (1996), 47-60.
- [4] K.C.B. Roes en R.J.M.M. Does, Shewhart-type charts in nonstandard situations, Technometrics 37 (1995), 296-303.

*Dr A. Di Bucchianico
Vakgroep Besliskunde en Stochastiek
Technische Universiteit Eindhoven*

Fisher, L.D., van Belle, G.

Biostatistics. A methodology for the health sciences

1993, Wiley series in probability and mathematical statistics: applied probability and statistics section,

John Wiley & Sons Ltd., Chichester,

xxii+991 blz, £ 62.00, ISBN 0-471-58465-7.

Dit boek is lijvig. Het weegt ruim anderhalf kilo. Voor het raadplegen is een plaats aan tafel aan te raden. Het boek is bedoeld voor mensen die op het biomedische terrein werken. Epidemiologen en tandartsen, biologen en biostatistici, gezondheidswetenschappers, verplegers en artsen (m/v).

De stof wordt gepresenteerd in 18 hoofdstukken. Te veel om ze allemaal op te noemen, maar de meeste in biomedische context gebruikte statistische methoden komen aan bod. Log-lineaire modellen voor kruistabellen, permutatie en randomizatie-toetsen, robuuste methoden, ANOVA en meervoudige regressie, multiple comparisons, discriminatie en classificatie, waaronder logistische regressie, principale componenten en factor analyse, survival analyse, ROC-curven. In het laatste hoofdstuk, Personal Postscript, beschrijven de auteurs enkele toepassingen die zij zelf als practiserende statistici als spannend, interessant en voldoening gevend hadden ervaren. Een sympathieke

poging om het enthousiasme voor het vakgebied over te brengen en de collaboratieve rol van de statistiek te benadrukken.

De presentatie is vrij overzichtelijk. Naast indeling in de paragrafen worden ook Definitions, Results, Facts, Examples en Notes apart aangegeven en genummerd. De auteurs schuwen de formules bepaald niet. Dit had iets minder gekund, bijvoorbeeld voor wat betreft de kwadraatsommen voor diverse ANOVA's. Er worden vrij veel details gegeven, wat ook tot uiting komt in de selectie van de in het boek opgenomen tabellen. Onder andere cumulatieve binomiale kansen en tabellen van kritieke waarden voor de Fishers exacte toets. De aanwezigheid van de door line-printer geproduceerde grafieken (gelukkig niet allemaal) geeft een enigszins gedateerde indruk. Maar dit zijn details. Het geheel overziend is mijn oordeel positief. Het boek bevat een indrukwekkende hoeveelheid van praktische voorbeelden, volgens de omslag meer dan 130, en ruim 390 opgaven. Deze zijn voorzien van uitgebreide achtergrondinformatie en verwijzingen naar originele publicaties. Antwoorden van de opgaven bevat het boek helaas niet.

Het boek wordt geacht toegankelijk te zijn ook voor mensen zonder enige statistische vooropleiding en met een vrij minimale kennis van wiskunde. Ik zou het een beginner echter niet aanraden. Deze is mijns inziens meer gebaat met een inleidende boek dat compacter is, beter aansluit bij de eigentijdse rekenfaciliteiten, de antwoorden op opgaven bevat, en een diskette met data als een bijlage heeft. Een voorbeeld van een boek dat aan deze wensen voldoet, op de laatste na, is Altman (1995).

Het sterkste aan het boek van Fisher en Van Belle vind ik de goed geselecteerde en gedocumenteerde praktijkvoorbeelden. Dit aspect maakt het aantrekkelijk voor docenten. De compleetheid maakt het goed bruikbaar als naslagwerk. Voor zelfstudie is het echter minder geschikt.

Verwijzingen:

[1] Altman D.G. (1995). *Practical Statistics for Medical Research*. Chapman & Hall, London.

dr V. Fidler

*Vakgroep Medische Statistiek
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.

Otten, G.D.

Statistical test limits in quality control

1996, CWI tracts,

Centrum voor Wiskunde en Informatica, Amsterdam,

143 blz, fl 35.00, ISBN 90-6196-469-5.

Maas, W.C.A.

Nonlinear $\mathcal{H}_\infty$ control: the singular case

1996, CWI tracts,

Centrum voor Wiskunde en Informatica, Amsterdam,

197 blz, fl 40.00, ISBN 90-6196-468-7.

Chow, S.L.

Statistical signifance. Rationale, Validity and Utility

1996, Sage, Thousand Oaks, California,

xi + 205 blz, £ 13.95, ISBN 0-7619-5205-5.

Duflo, M.

Stochastic recursive methods

1997, Applications of mathematics,

Springer-Verlag, Berlin,

385 blz, DM 128.00, ISBN 3-540-57100-0.

Linden, W.J. van der, Hambleton, R.K. (Eds.)

Handbook of modern item response theory

1996, Springer-Verlag, Berlin,

blz, DM 108.00, ISBN 0-387-94661-6.

Kimura, H.

Chain-scattering approach to H-infinity-control

1995, Systems and control: foundations and applications,
Birkhäuser, Basel,
250 blz, DM 94.00, ISBN 3-7643-3787.

Bardi, M., Dolcetta, I.C.

Optimal control & viscosity solutions of Hamilton-Jacobi-Bellman equations

1996, Progress in systems and control,
Birkhäuser, Basel,
500 blz, DM 178.00, ISBN 3-7643-3640-4.

Peters, M.A., Iglesias, P.

Minimum entropy control for time-varying systems

1996, Systems and control: foundations and applications,
Birkhäuser, Basel,
200 blz, DM 94.00, ISBN 3-7643-3972-1.

Chauvin, B., Cohen, S., Rouault, A.

Trees. Workshop in Versailles, June 14–16, 1995

1996, Progress in probability,
Birkhäuser, Basel,
158 blz, DM 88.00, ISBN 3-7643-5453-4.

Borg, I., Groenen, P.

Modern multidimensional scaling. Theory and applications

1996, Springer series in statistics,
Springer-Verlag, Berlin,
480 blz, DM 88.00, ISBN 0-387-94845-7.

Reyment, R., Jøreskog, K.G.

Applied Factor Analysis in the Natural Sciences

1997, Cambridge University Press, Cambridge,
xii+371 blz, £ 55.00, ISBN 0-521-41242-0.

Long, J.S.

Regression models for categorical and limited dependent variables

1997, Sage, Thousand Oaks, California,

330 blz, £ 37.00, ISBN 0-8039-7374-8.

di Bucchianico, A.

Probabilistic and analytical aspects of the Umbral calculus

1997, CWI Tract 119,

Centrum voor Wiskunde en Informatica, Amsterdam,

148 blz, fl 35.00, ISBN 90-6196-471-7.

Bomze, I.M., Csendes, T., Horst, R., Pardalos, P.M. (ed.)

Developments in global optimization

1997, Nonconvex optimization and its applications **18**,

Kluwer, Dordrecht,

360 blz, fl 240.00, ISBN 0-7923-4351-4.

Ansari, N., Hou, E.

Computational intelligence for optimization

1997, Kluwer, Dordrecht,

240 blz, fl 205.00, ISBN 0-7923-9838-6.

Murli, A., Toraldo, G.

Computational issues in high performance software for nonlinear optimization

1997, Kluwer, Dordrecht,

192 blz, fl 215.00, ISBN 0-7923-9862-9.

West, M., Harrison, J.

Bayesian forecasting and dynamic models

1997, Springer series in statistics,

Springer-Verlag, Berlin,

725 blz, DM 88.00, ISBN 0-387-94725-6.

Lehmann, E.L.

Testing statistical hypotheses

1997, Springer texts in statistics,
Springer-Verlag, Berlin,
625 blz, DM 98.00, ISBN 0-387-94919-4.

Lederer, P.J., Karmarkar, U.S. (ed.)

The practice of quality management

1997, Kluwer, Dordrecht,
328 blz, fl 230.00, ISBN 0-7923-9864-5.

Pallaschke, D., Rolewicz, S.

Foundations of mathematical optimization. Convex analysis without linearity

1997, Mathematics and its applications,
Kluwer, Dordrecht,
596 blz, fl 395.00, ISBN 0-7923-4424-3.

Jaiswal, N.K.

Military Operations Research

1997, International series in operations research and management science,
Kluwer, Dordrecht,
400 blz, fl 260.00, ISBN 0-7923-9858-0.

Jansen, B.

Interior point techniques in optimization. Complementarity, sensitivity and algorithms

1997, Applied optimization,
Kluwer, Dordrecht,
292 blz, fl 210.00, ISBN 0-7923-4430-8.

Hayakawa, T., Aoshima, M., Shimizu, K. (ed.)

Multivariate statistical analysis in honor of Professor Minoru Siotani on his 70th birthday. Volume II of the proceedings of MSI-2000

1997, American Sciences Press, Syracuse, New York,
266 blz, \$ 125.00, ISBN 0-935950-39-7.

Sachs, L.

Angewandte Statistik. Anwendung statistischer Methoden

1997, Springer-Verlag, Berlin,

884 blz, DM 98.00, ISBN 3-540-60494-4.

Firebaugh, G.

Analyzing repeated surveys

1997, Quantitative applications in the social sciences,

Sage, Thousand Oaks, California,

vii+72 blz, £ 7.95, ISBN 0-8039-7398-5.

Oproep voor boekbesprekers

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

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

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

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

