

## Boekbesprekingen

### *Redactie*

Arend Oosterhoorn  
 Van Doorne's Transmissie bv  
 Postbus 500  
 5000 AH TILBURG

telefoon 013 - 640319  
 fax 013 - 636590

## Besproken boeken in Kwantitatieve Methoden nr. 41

|                                                                                                            |     |
|------------------------------------------------------------------------------------------------------------|-----|
| <i>Thomas R. Fleming and David P. Harrington</i><br>Counting Processes and Survival Analysis               | 162 |
| <i>C. Chatfield</i><br>The analysis of time series. An introduction. (fourth edition).                     | 163 |
| <i>P.J.M. Verschuren.</i><br>Structurele modellen tussen theorie en praktijk                               | 165 |
| <i>Noel A.C. Cressie</i><br>Statistics for spatial data                                                    | 167 |
| <i>Richard A. Johnson &amp; Gouri K. Bhattacharyya</i><br>Statistics, Principles and Methods (2nd edition) | 170 |
| <i>Morgenthaler, Stephan &amp; Tukey, John W. (eds.).</i><br>Configural polysampling                       | 172 |
| <i>Ashish Sen en Muni Srivastava</i><br>Regression analysis: theory, methods and applications              | 173 |
| <i>Joep Brinkman</i><br>Cijfers spreken, statistiek en methodologie voor het hoger onderwijs               | 175 |

|                                                                                                                          |     |
|--------------------------------------------------------------------------------------------------------------------------|-----|
| <i>Prem S. Mann</i><br>Introductory Statistics                                                                           | 176 |
| <i>T.J. Hastie &amp; R.J. Tibshirani</i><br>Generalized Additive Models                                                  | 178 |
| <i>T.H.B. Hendriks en P. van Beek</i><br>Optimaliseringstechnieken: principes en toepassingen (3e druk)                  | 180 |
| <i>Geoffrey J. McLachlan</i><br>Discriminant analysis and statistical pattern recognition                                | 181 |
| <i>Carl-Erik Särndal, Bengt Swensson, Jan Wretman</i><br>Model Assisted Survey Sampling/Springer Series in Statistics    | 183 |
| <i>A.S. Hedayat &amp; Bikas K. Sinha</i><br>Design and inference in finite population sampling                           | 184 |
| <i>J. Muilwijk, T.A.B. Snijders en J.J.A. Moors</i><br>Kanssteekproeven                                                  | 185 |
| <i>A.I. Dale</i><br>A history of inverse probability, from T. Bayes to K. Pearson                                        | 186 |
| <i>M.J. Crowder, A.C. Kimber, R.L. Smith en T.J. Sweeting.</i><br>Statistical analysis of reliability data.              | 190 |
| <i>E.K. Foreman</i><br>Survey sampling principles                                                                        | 192 |
| <i>Alan Pankratz</i><br>Forecasting with dynamic regression models                                                       | 193 |
| <i>T.P. Hutchinson</i><br>Ability, partial information, guessing: statistical modelling applied to multiple-choice tests | 195 |
| <i>T.P. Hutchinson</i><br>Controversies in Item Response Theory                                                          | 197 |
| <i>Lyman Ott and William Mendenhall</i><br>Understanding statistics                                                      | 198 |

**Thomas R. Fleming and David P. Harrington**

*Counting Processes and Survival Analysis*

John Wiley & Sons, New York, 1991, xiv + 429 pag., ISBN 0-471-52218-X, £ 39.80.

In de laatste 25 jaar heeft de ontwikkeling van statistische methoden voor het analyseren van gecensureerde gegevens een hoge vlucht genomen. Vele nieuwe methoden werden geïntroduceerd, waarvan de belangrijkste het proportional hazards regressie model van Cox was. Aanvankelijk hadden deze methoden nogal een ad hoc karakter, maar in het laatste decennium is veel aandacht besteed aan het wiskundig onderbouwen van deze methoden. Deze wiskundige onderbouwing steunt vooral op de theorie van de telprocessen, martingalen en stochastische integralen. Dit heeft, behalve tot een unificatie van de theorie en tot nieuwe inzichten in de bestaande methoden, ook geleid tot belangrijke generalisaties en nieuwe toepassingsmogelijkheden.

Dit boek is bij mijn weten het eerste algemeen gerichte survival analyse boek waarin de analyse van gecensureerde gegevens op deze moderne manier wordt benaderd. Als voordelen van deze aanpak noemen de schrijvers:

- a. Unificatie van de theorie; het geeft een algemeen wiskundig kader om relatief gemakkelijk de eigenschappen van allerlei methoden te bestuderen.
- b. Generalisaties naar ingewikkelder situaties, zoals meer events per individu of meer typen events, zijn relatief gemakkelijk.
- c. Deze benadering voorziet in een goed denkkader voor ingewikkelde toepassingen.

Een nadeel is dat de lezer zich eerst door een flink brok (tamelijk moeilijke) wiskunde moeten heen worstelen. De benodigde theorie van telprocessen en martingalen wordt in het boek zelf tamelijk grondig behandeld. Een groot gedeelte van het boek is daar dan ook aan gewijd, zoals blijkt uit de volgende inhoudsopgave:

1. The counting process and martingale framework.
2. Local square integrable martingales.
3. Finite sample moments and large sample consistency of tests and estimators.
4. Censored data regression models and their application.
5. Martingale central limit theorem.
6. Large sample results of the Kaplan-Meire estimator.
7. Weighted log rank statistics.
8. Distribution theory for proportional hazards regression.

Voor degenen die zich de moderne theorie van de survival analyse willen eigen maken en niet terugschrikken voor een behoorlijke portie wiskunde, lijkt mij dit boek een aanwinst. Resultaten die tot nu verspreid in de literatuur gezocht moesten worden, zijn nu geïntegreerd en compact bijeengebracht en helder beschreven. Ik vrees echter dat voor de meeste toegepaste statistici dit boek vanwege het hoge wiskundige nivo tamelijk ontoegankelijk zal blijken te zijn.

Th. Stijnen

Instituut voor Epidemiologie en Biostatistiek, EUR.



### C. Chatfield

*The analysis of time series. An introduction. (fourth edition).*

Chapman and Hall, Londen, 1989, xiii + 241 pag., ISBN 0-412-31820-2, £ 13.50.

De vierde editie van Chatfields "The analysis of Time Series" is een bijna geheel herziene versie van de vorige edities. De eerste negen hoofdstukken zijn grotendeels herschreven en hoofdstuk 10 is geheel nieuw en behandelt "state-space models". Hoofdstuk 11 is ook geheel herzien; hierin worden verschillende losse onderwerpen behandeld. Het volgende overzicht geeft aan dat een groot aantal onderwerpen m.b.t. tijdreeks-analyse aan de orde komen.

In hoofdstuk 1 worden verschillende soorten tijdreeksen, de doelstellingen en benaderingen van de tijdreeks-analyse gepresenteerd, alsmede een kort overzicht van verdere literatuur. Hoofdstuk 2 bevat een goed overzicht van eenvoudige beschrijvende technieken, data transformatie en analyses van tijdreeksen met een trend en/of seizoensfluctuaties d.m.v. verschillende filters. Diverse kansmodellen voor tijdreeksen met hun bijbehorende autocorrelatiefunctie worden in hoofdstuk 3 besproken.

In hoofdstuk 4 komen schattingsprocedures voor het tijddomein aan de orde: het schatten en interpreteren van de autocorrelatiefunctie, het fitten van een AR en van een MA proces met het bepalen van hun orde alsmede ARIMA en SARIMA modellen. Het hoofdstuk wordt afgesloten met de bespreking van de analyse van de residuen en met algemene opmerkingen over het maken van een model (specificatie, fitten en verificatie).

Verschillende voorspellingstechnieken worden in hoofdstuk 5 geïntroduceerd en met elkaar vergeleken. Verder wordt er in het kort de Box-Jenkins methode besproken en tenslotte wordt het hoofdstuk met enige duidelijke voorbeelden afgesloten.

Hoofdstuk 6 gaat over stationaire processen in het frekwentie domein. Hierin komen o.a. aan de orde: de spectrum verdelingsfunctie alsmede de afleidingen voor de spectraal dichtheidsfuncties van verschillende stationaire processen.

Analoog aan hoofdstuk 4 wordt er in hoofdstuk 7 ingegaan op de schattingsprocedures voor de spectraal analyse met behulp van onder andere periodogrammen.

In hoofdstuk 8 worden bivariate processen besproken. De nadruk ligt er op de correlatie tussen twee tijdreeksen. De definitie van de cross-variantie in de cross-correlatie functies voor tijdreeksen in het tijddomein worden gegeven alsmede de theorie voor het schatten en interpreteren ervan. Voor tijdreeksen in het frekwentie domein wordt op analoge wijze de cross-spectraal functie besproken.

In hoofdstuk 9 worden lineaire systemen met één input en één output variabele beschouwd voor zowel tijdreeksen in het tijds- als in het frekwentie-domein. M.b.t. de indentifikatie van een lineair systeem wordt er nader op het schatten van de frekwentie-response functie ingegaan en op het schatten van het tranferfunctie model van Box-Jenkins. Tenslotte worden er systemen besproken waarbij feedback een rol speelt.



Hoofdstuk 10 is een nieuw hoofdstuk (ten opzichte van eerdere edities). Er wordt op heldere wijze ingegaan op de state-space modellen (hoe onder andere AR-modellen als een state-space model te schrijven zijn) en op het Kalman filter.

Tenslotte bevat hoofdstuk 11 een aantal losse onderwerpen zoals: control theorie, niet-stationaire tijdreeksen, niet-lineaire modellen, middelen voor model-identificatie (o.a. Akaike's final predictor error, Akaike's information criterium, Parzen's autoregressive transferfunction criterium) en multivariate modellen.

In vier appendices worden behandeld:

- a. Fourier, Laplace en Z-transformaties
- b. Dirac's delta functie
- c. Covariantie
- d. enkele uitgewerkte voorbeelden met behulp van het statistische pakket Minitab.

Het boek is uiterst prettig leesbaar en is zoals de ondertitel belooft, een echte inleiding in de tijdreeks-analyse. Op heldere wijze weet Chatfield met een minimum aan formules de basisgedachten van de tijdreeks-analyse over te brengen. Daar waar hij te summier is en mathematische moeilijkheden niet vermeld vanwege de overzichtelijkheid, geeft hij dit zelf aan en verwijst er naar recente en overzichtelijke literatuur. De referenties (172 in totaal) zijn bijgewerkt wat het boek actueel maakt.

De hoofdstukken 2 t/m 10 worden afgesloten met enkele vragen, waarvan de antwoorden achter het boek vermeld staan.

Ook niet wiskundigen, die zich in tijdreeksen willen verdiepen zullen dit boek met veel plezier lezen. Je zou het de tegenhanger van het recente boek van Brockwell and Davis "Time Series: Theory and Methods" kunnen noemen, waarin juist op degelijke wijze op de wiskundige achtergronden en rekentechnieken in tijdreeks-analyse wordt ingegaan.

Beide boeken vullen elkaar in didactische zin dan ook uitstekend aan.

#### Referentie

P.J. Brockwell and R.A. Davis, 1991. *Time series: Theory and Methodes*. Springer Verlag: Berlin Heidelberg New York (ISBN 3-530-97429-6)

Arnold Dekkers

RIVM, Bilthoven

**P.J.M. Verschuren.**

*Structurele modellen tussen theorie en praktijk*

Het Spectrum bv, Utrecht, 1991, 662 pag., ISBN 90-274-2579-5, f 79.90.

Dit lijvige boek heeft tot doel de lezer vanaf de eerste beginselen van de covariantie structuur analyse mee te nemen naar de meest recente methoden en programmatuur. Het is bestaat uit 18 hoofdstukken verdeeld over 5 delen. Dat weinig voorkennis wordt verondersteld, blijkt uit de zeer elementaire behandeling van zaken als bivariate regressie, begrippen als covariantie, correlatie en standaarddeviatie als ook uit de aanwezigheid van appendices waarin elementaire achtergrond kennis uit de doeken wordt gedaan zoals het begrip kansverdeling, regels voor sommatie en zelfs elementaire wiskundige operaties (e.g.,  $-[-a] = a$ ). Er worden veel onderwerpen behandeld die bij de meeste studenten in de sociale wetenschappen bekend verondersteld mogen worden.

Een groot deel van het boek gaat uit van structurele modellen waarvoor de te schatten parameters een expliciete functie van de observaties vormen, dus modellen waarvan de parameters in gesloten vorm opgelost kunnen worden. Hierdoor komt de nadruk te liggen op recursieve structurele modellen met gemeten variabelen en ongecorrleerde storingstermen. Later in het boek worden andere modellen behandeld waarvoor iteratieve schattingsprocedures nodig zijn. Een nadeel hiervan is enerzijds dat een aantal onderwerpen twee keer wordt behandeld, waaronder verlies functies (maximum likelihood, unweighted least squares), toetsing, en bepaalde aspecten van model evaluatie. Anderzijds worden door de tekst heen wel de complicaties genoemd die ontstaan bij meer algemene modellen bijvoorbeeld voor de schattingsprocedure. Dit geeft een wat ongeorganiseerde indruk. Samenvattingen aan het einde van de hoofdstukken ontbreken. Dit is met name bij een boek van deze omvang een gemis.

Noodzakelijkerwijs worden in een inleidend boek bepaalde zaken benadrukt ten koste van andere. Echter de afwezigheid van een aantal belangrijke onderwerpen is opvallend en betekent in ieder geval dat de op de achterflap gedrukte belofte de lezer mee te nemen naar de nieuwste methoden en programmatuur niet wordt waargemaakt. LISREL 6, bijvoorbeeld, wordt consequent gebruikt terwijl LISREL 7 sedert 1988 verkrijgbaar is. Dit heeft tot gevolg dat 'normalised residuals' gerapporteerd worden terwijl het bekend is dat deze incorrect zijn. 'Standardised residuals' uit LISREL 7 verdienen de voorkeur (Jöreskog en Sörbom, 1989). Browne's (1984!) "asymptotically distribution-free" verlies functie en de implementatie hiervan in LISREL 7 (weighted least squares, WLS, zie Jöreskog en Sörbom, 1989) wordt niet behandeld. De WLS implementatie is met name in praktisch opzicht belangrijk omdat deze geschikt is voor de analyse van niet normaal verdeelde continue variabelen als ook discreet verdeelde data samengevat in tetrachorische of polychorische correlatie matrices. Het behoeft geen betoog dat in praktisch sociaal wetenschappelijk onderzoek discrete variabelen eerder regel dan uitzondering zijn. Opvallend is dan ook dat in dit boek het onderscheid tussen dichotome en polytome data wordt genoemd (blz. 45, 77) zonder dat men zich bekommert over het aantal antwoordcategorieën in het polytome geval. Afwezig is tevens een behandeling van passingsindices. De indices AGFI, GFI, chi-kwadraat toetsingsgrootheid en de ratio van chi-kwadraat (chi) tot de



vrijheidsgraden (vg),  $\chi^2/vg$ , worden genoemd, maar andere in de literatuur veelvuldig voorkomende indices zijn afwezig: Bentler en Bonnett index, Tucker-Lewis index, Akaike's informatie criterium, Cudeck en Browne's kruisvalidatie index, en de op geneste modellen van toepassing zijnde toetsingsgrootheden (zie Bollen, 1989).

Tot de laatst genoemde behoort de bekende hierarchische  $\chi^2$ -kwadraat toets, een boven de t-toets te verkiezen procedure om de significantie van individuele parameters te bepalen (zie Neale, et al. 1989). Statistisch onderscheidingsvermogen (zie bijvoorbeeld Satorra, Saris en de Pijper, 1991), wederom een onderwerp van groot praktisch belang bij hierarchische toetsingen, wordt niet behandeld. Wel wordt nogal arbitrair advies gegeven:  $N > 100$  (blz. 249, 465) enerzijds,  $N = 40$  à  $50$  anderzijds (blz. 24, 267). Ten aanzien van de kritische waarden van AGFI en  $\chi^2/vg$  wordt de lezer eveneens voorzien van arbitrair advies, waarbij de richtlijn dat  $\chi^2/vg > 3$  duidt op een onvoldoende passing niet te onderschrijven is wanneer men een model toetst met weinig vrijheidsgraden. Van AGFI en GFI wordt overigens beweerd dat ze beter bestand zijn tegen afwijkingen van normaliteit dan de  $\chi^2$ -kwadraat toets en dat ze onafhankelijk zijn van de steekproef grootte ( $N$ ). Het laatst genoemde is onjuist, zoals blijkt uit simulatie onderzoek van Marsh, Balla en McDonald (1988). Bovendien lijkt robuustheid gegeven afwijkingen van normaliteit mij niet van toepassing op informele passingsmaten met onbekende verdelingen. Van de  $\chi^2$ -kwadraat toets wordt gezegd dat het, gezien de voorwaarden die eraan gesteld worden, raadzaam is om deze op te vatten als een informele indicator in plaats van een formele toetsingsgrootte (blz. 469). Dit is een redelijk advies, maar het strookt niet met het gegeven advies de standaardfout van een geschatte parameter te gebruiken om hypothesen te toetsen ten aanzien van de waarden van de schatting (blz. 466). Immers, als de  $\chi^2$ -kwadraat toets niet te vertrouwen is, dan zijn de standaardfouten dat ook niet.

Ik tel betrekkelijk weinig drukfouten. Storend is echter wel dat er iets mis is gegaan met de referenties. Ik tel tenminste 20 in de tekst voorkomende referenties die niet, of niet correct, voorkomen in de literatuur lijst. De auteur spaart zichzelf daarbij overigens niet: bij de referentie Potman en Verschuren staat als jaartal 1989 in de tekst (blz. 167) en 1987 in literatuur lijst (blz. 659).

Het belangrijkste bezwaar tegen dit boek is dat het onduidelijk blijft voor wie het is geschreven. Voor wie niets weet van statistiek, zijn er betere inleidende Nederlandse teksten (bijvoorbeeld de delen van Van de Brink en Koele, 1986). Het is te elementair en te traag voor iemand die opgelet heeft tijdens zijn/haar methodologie colleges. Als 'textbook' is het onge-schikt omdat te veel relevante informatie ontbreekt.

#### Referenties.

- Satorra, A., Saris, W.E. en de Pijper, W.M. (1991). A comparison of several approximations to the power function of the likelihood ratio test in covariance structure analysis. *Statistica Neerlandica*, 45, 2, 173-185.
- Jöreskog, K.G. en Sörbom, D. (1989). *LISREL7, A guide to the program and applications*. Chicago: SPSS Inc.

- Marsh, H.W., Balla, J.R. en McDonald, R.P. (1988). Goodness of fit indexes in confirmatory factor analysis: the effect of sample size. *Psychological Bulletin*, 100, 107-120.
- Neale, M.C., Heath, A.C., Hewitt, J.K., Eaves, L.J. and Fulker, D.W. (1989). Fitting genetic models with LISREL: hypothesis testing. *Behavior Genetic*, 12, 37-49.
- Van den Brink, W.P. en Koele, P. (1986). *Statistiek, deel I, II, en III*. Amsterdam: Boom Meppel.
- Bollen, K.A. (1989). *Structural equations with latent variables*. New York: John Wiley and Sons.
- Browne, M.W. (1984) Asymptotically distributions-free methods for the analysis of covariance structures. *British Journal of Mathematical and Statistical Psychology*, 37, 62-83.

Conor V. Dolan

Faculteit der Psychologie, Universiteit van Amsterdam.

**Noel A.C. Cressie**

*Statistics for spatial data*

John Wiley & Sons, New York, 1991, xx + 900 pag., ISBN 0-471-84336-9, £ 71.00 .

Wat heeft exploratie-onderzoek ten behoeve van een kolenmijn te maken met het verschijnsel van wiegedood? Wat heeft dat verschijnsel uitstaande met de ruimtelijke verdeling van naaldbomen in een bos? Kan er enig verband bestaan tussen die ruimtelijke verdeling en de groei van tumoren, en hoe verhoudt zich dit laatste weer tot het vermelde exploratie-onderzoek? Een antwoord op deze vragen vindt men in dit boek over ruimtelijke statistiek.

Cressie brengt deze problemen met elkaar in verband door de uniforme beschrijving van een multivariaat random proces  $\{Z(s) \mid s \in D \subset \mathbb{R}^n\}$ . Voor geostatistische data (exploratie onderzoek) is  $D$  een vaste deelverzameling van en  $Z(s)$  is een random vector te  $s \in D$ . Als  $D$  een vaste verzameling is van (ir-)regulier verdeelde punten, dan betreft het roosterdata (wiegedood per district). Ruimtelijke patronen worden beschreven als  $D$  zelf een puntproces is: is  $Z(s)$  een random vector dan is sprake van (gemarkeerde) puntpatronen (verdeling van naaldbomen), en is  $Z(s)$  een random verzameling dan worden objecten beschreven (tumoren).

Doelstelling van de auteur is de presentatie van statistiek voor ruimtelijke data aan zuivere en toegepaste wetenschappers. Dat houdt in dat naast theoretische aspecten van de statistische modellering ook praktische aspecten aan de orde komen. De auteur heeft niet geschroomd ruime aandacht te schenken aan het laatste aspect. Hij onderstreept het belang van data analyse voorafgaande aan de modellering en toont dat in de praktijk aan. Inderdaad, volgens Cressie zijn data de wortels van de statistische analyse. Daarom heeft hij van een aantal in het boek behandelde voorbeelden de data opgenomen, zodat de lezer zelf aan de slag kan zonder zich te



hoeven bekommeren om een moeizame speurtocht daarnaar. Na de data analyse volgen de statistische modellering en de schatting van modelparameters.

Een groot aantal hoofdstukken staat inhoudelijk tussen de toepassing en het zware theoretische geschut in. Deze zijn zowel voor de praktische gebruiker als voor de zuivere wetenschapper van belang en leggen zodoende een link tussen beiden. Ze behandelen de statistische modellen en hun eigenschappen, zoals robuustheid ten opzichte van de afwijking van de verdeling van data van het vooronderstelde model. Alle modellen zijn geparameteriseerd en dus komt ook de parameterschatting aan de orde: zuiverheid of onzuiverheid, efficiency, resistentie ten opzichte van uitschieters en de (schatting van de) variantie in de geschatte parameter.

In de meest theoretische hoofdstukken wordt de hier vermelde stof verder verdiept en uitgebreid. Zo komt men informeel te weten dat bij een compleet willekeurige verdeling van punten op een vlak, deze verdeeld zijn volgens een Poisson proces waarbij de parameter voor de dichtheid van punten per oppervlakte een constante is. Bij een inhomogeen Poisson proces is deze parameter een functie van de plaats. Het kan ook anders: een inhomogeen Poisson proces wordt beschreven door middel van een willekeurige Radon-maat en is slechts dan homogeen als die maat evenredig is met de Lebesgue-maat.

Overigens komt men herhaaldelijk verwijzingen tegen in de meer toegepaste hoofdstukken naar theoretische en *vice versa*. Ook hier lijkt Cressie te willen streven naar samenhang. Bij de statistische modellering dient men niet blindweg theorie in praktijk toe te passen, maar dient men goed op de hoogte te zijn van de theoretische achtergronden en implicaties van een model. Anderszijds is verdergaande theorievorming gediend met een goed begrip van de moeilijkheden die men in de praktische toepassing tegenkomt en van de soort problemen waarmee de praktijk geconfronteerd wordt. Deze laatste diversiteit komt duidelijk tot uiting.

Het boek is verdeeld over drie delen: geostatistische data, roosterdata en ruimtelijke patronen. Maar in het hoofdstuk over toepassingen van de geostatistiek, betreffende grondwaterstroming, de potentiaal van onverzadigd bodemvocht, bodemvocht infiltratie, wiegedood, graanoogst (een probleem stammend uit 1911 (!)) en zure regen, horen de onderwerpen over wiegedood en graanoogst feitelijk in het tweede deel thuis waar zij ook nogmaals behandeld worden. Evenzo blijkt in het derde deel, bij het probleem van de verdeling van naaldbomen, dat na datareductie door middel van tellingen per eenheidsvierkant een behandeling volgens de discrete roosterdata methode of de continue geostatistische ook tot de mogelijkheden behoort. Dit benadrukt des te meer de unificatie die de auteur beoogt.

Het eerste deel over geostatistiek behandelt mogelijkheden van data-analyse, waarbij de median polish (mediaanschaaf?) methode, ook met betrekking tot het vervolg, als resistente methode veel aandacht krijgt. Daarna, na het verwijderen of editen van uitschieters, komt het schatten van variogrammen en het fitten van geschikte variogrammodellen aan de orde. Simpele, gewone, robuuste, universele, of median polish kriging (optimum interpolatie voor o.a. tijdreeksanalisten en meteorologen) is de methode om data te voorspellen. Het al vermelde hoofdstuk met toepassingen komt daarna, waarna nog een hoofdstuk met specialistische onderwerpen volgt.

Het tweede deel over roosterdata behandelt conditioneel gespecificeerde ruimtelijke (Gaussische) modellen naast simultaan gespecificeerde. Conditioneel gespecificeerde modellen geven de kansverdeling in een roosterpunt gegeven een conglomeraat van buurpunten. Simultaan gespecificeerde modellen geven de kansverdeling van een conglomeraat van roosterpunten. Deze modellen worden gebruikt in samenhang met discrete waarden voor data (0-1 data of tellingen) en met continue waarden. Parameter schatters en hun eigenschappen evenals enige praktische voorbeelden worden behandeld. Bij deze praktische voorbeelden komen telkens 2-dimensionale aangrenzende gebieden aan de orde, waarbij het rooster gecreëerd wordt door een gebied een representatieve coördinaat toe te kennen. Enerzijds wordt hier veronachtzaamd dat het feitelijk data over oppervlakken betreft en geen data in punten, anderzijds ontstaat het probleem van het toekennen van een geschikte representant. Bijvoorbeeld in het geval van het wiegedood probleem kiest Cressie de hoofdstad van een bestuurlijk district, waar een ander het zwaartepunt van dat district zou kiezen en nog weer een ander het zwaartepunt gewogen met de bevolkingsdichtheid.

Het derde deel behandelt ruimtelijke patronen. Het bevat relatief de meeste theoretische hoofdstukken. De voornaamste opgave bij pure ruimtelijke puntprocessen is te ontdekken of er sprake is van volledige willekeur (geen interactie), clustering (attractieve interactie) of regelmatigheid (repulsieve interactie). De behandeling van de rand van een gebied blijkt van groot belang te zijn in verband met naaste-buur detectie. Behalve de pure puntprocessen worden ook gemarkeerde processen beschouwd. Aan een gebeurtenis worden dan eigenschappen, zoals bijvoorbeeld kleur of grootte, toegekend. Objecten zijn bijvoorbeeld lager dimensionale representaties van een verschijnsel in hogere dimensies, zoals een-dimensionale doorsneden van een drie-dimensionaal object in geval van tomografie of twee-dimensionale doorsneden in geval van Röntgen fotografie. Bij tumoren is de plaats van de kern van een woekering het puntproces  $D$  en de woekering zelf de random verzameling  $Z(s)$ . Het Boolean model wordt hier als statistisch model geïntroduceerd, uitgewerkt en toegepast.

Het boek geeft de state-of-the-art van de ruimtelijke statistiek. Het is een samenvatting van al bekende zaken en bevat daarnaast hoofdstukken die grensverleggend zijn. Aan statistiek voor data in ruimte en tijd wordt getipt, maar behandeling van dit soort problemen wordt uitgesteld voor de toekomst. Cressie geeft uitgebreide referenties: 70 pagina's. De index is gesplitst in een over auteurs en een over onderwerpen. Het kan, onder begeleiding, als leerboek gebruikt worden in het na-kandidaats curriculum, waarbij de docent zelf opgaven zal moeten formuleren, omdat deze ontbreken. Hoofdvak studenten moeten in staat geacht worden zelf met vrucht enige capita te bestuderen.

Ik denk dat dit boek voor allen die zich beroepsmatig met ruimtelijke statistiek bezighouden, theoretisch of toegepast, een uitstekende bron van samengevatte bekende en nieuwe informatie is. Het mag zeker in geen enkele zichzelf respecterende bibliotheek van een (toegepast) wetenschappelijk instituut waar statistiek bedreven wordt ontbreken. De prijs bedraagt ongeveer 210 Nederlandse guldens, wat in absolute zin een behoorlijk bedrag is. Evenwel kan men ook



argumenteren dat qua omvang en qua inhoud dit ene boek wel drie boeken waard is, en dan is 70 gulden per deel een zeer schappelijke prijs.

Jan v. Eijkeren  
RIVM, Bilthoven

**Richard A. Johnson & Gouri K. Bhattacharyya**

*Statistics, Principles and Methods (2nd edition)*

John Wiley & Sons, New York, 1992, xvii + 686 pag., ISBN 0-471-54842-1, £ 19.50.

Dit boek is bestemd voor studenten, die in een korte tijd kennis moeten maken met de grondbeginselen van de statistiek. De nadruk ligt op de toepassingen van de diverse technieken.

Het boek bevat een hoofdstuk over het beschrijven van bivariate gegevens, dat volgt op het gebruikelijke hoofdstuk over beschrijvende statistiek. Studenten raken daardoor al snel vertrouwd met de belangrijkste terminologie en de samenhang tussen de hoofdstukken die nog gaan komen. Het toetsen van hypothesen wordt twee keer behandeld. Zowel op de gebruikelijke manier als direct na de introductie van de binomiale verdeling. Dit is zeer verhelderend, temeer omdat bij VWO wiskunde A het toetsen ook op dergelijke wijze geïntroduceerd werd. Tevens bevat het boek een aantal paragrafen over het nagaan van modelveronderstellingen en transformaties.

De opgaven in het boek gaan over de zeer vele toepassingsgebieden van de statistiek. Dit geeft een heel aardig beeld van de onbegrensde mogelijkheden van het vak. Verder kunnen enkele opgaven met behulp van MINITAB (of een ander pakket) opgelost worden. De uitleg bij de commando's is duidelijk. Ook in de verschillende paragrafen komt veel computeroutput voor.

Ieder hoofdstuk begint met een toepassing van datgene wat komen gaat. Daarna volgen 6 à 7 paragrafen, waarin theorie en voorbeelden elkaar afwisselen. Belangrijke opmerkingen en formules staan in een groen kader. Na elke para-graaf volgen opgaven. Aan het eind van ieder hoofdstuk staan de belangrijkste zaken nog eens op een rij. "Review exercises" sluiten een hoofdstuk af. In een appendix worden eigenschappen van de verwachting en variantie afgeleid.

De uitleg van de theorie is zeer duidelijk. Door de groen-zwarte lay out ziet het boek er verzorgd uit. De notatie is zorgvuldig en consequent.

Het boek bestaat uit 16 hoofdstukken. Globaal is de opbouw als volgt. Na het introductiehoofdstuk (H1) volgen twee hoofdstukken over beschrijvende statistiek (H2 en H3). De modus ontbreekt in de theorieparagraaf en het aantal behandelde spreidingsmaten is wat mager. Percentielen en EDA technieken komen ook aan bod. De boxplot wordt echter te eenvoudig behandeld. Dit is jammer omdat de SAS output aan het eind van het hoofdstuk wél de "Tukey versie" laat zien.

In het hoofdstuk over het beschrijven van bivariate gegevens komen kruistabellen, scatterplots, en de correlatiecoëfficiënt (en waarom men daar erg voorzichtig mee moet zijn) aan de orde. Zelfs regressie wordt al even bekeken. Het analyseren van kruistabellen en de echte regressie-analyse komen pas later.

Kansrekening (H4) en kansverdelingen (H5) zijn de volgende onderwerpen. Hoewel een aselechte steekproef wel gedefinieerd wordt, kunnen studenten er niet zelf een trekken omdat er geen aselechte getallen in het boek staan. Hoofdstuk 5 gaat over discrete stochasten. Ik mis de stelling:  $E[g(X)] = \sum g(x)f(x)$ . Nu komt  $\text{VAR}[X]$  enigszins uit de lucht vallen.

Jammer genoeg komt er maar één discrete verdeling aan bod, namelijk de binomiale verdeling (H6). Pas in de opgaven komt de Poisson verdeling naar voren. Ook het aantal continue verdelingen blijft beperkt tot één (de Normale verdeling H7). In het boek wordt de notatie  $N(\mu, \sigma)$  gehanteerd, terwijl  $N(\mu, \sigma^2)$  gangbaarder is.

Via H8, steekproefverdelingen, wordt de stap naar inferentie gezet. Dit is een zeer helder hoofdstuk met aardige voorbeelden.

Schattingstheorie en de Z-toets voor  $\mu$  en P bij een grote steekproef komen in H9 aan de orde. De T toets en de  $\chi^2$  toets voor  $\sigma$  komen in H10 ter sprake. Tevens wordt in dat hoofdstuk iets gezegd over de relatie tussen schatten en toetsen, en over robuustheid.

Hoofdstuk 11 gaat over het vergelijken van twee "behandelingen". Men gebruikt de notatie en terminologie die bij statistische proefopzetten gebruikelijk is. Er wordt aandacht besteed aan de situatie dat  $\sigma_1$  en  $\sigma_2$  ongelijk verondersteld worden. Ook via boxplots worden behandelingen met elkaar vergeleken.

De volgende twee hoofdstukken gaan over regressie-analyse.  $R^2$  wordt afgeleid en het algemene lineaire model wordt gepresenteerd. Verder is er aandacht voor transformaties en het nagaan van de modelveronderstellingen via residuën.

Ook het hoofdstuk over de analyse van categorische gegevens (H14) is zeer duidelijk geschreven. De analogie tussen het toetsen op homogeniteit in een 2x2 tabel en de P-toets voor twee populaties komt ook aan de orde.

ANOVA is het onderwerp van H15. Het ongebalanceerde volledig gerandomiseerde schema en het gerandomiseerde blokkenschema worden besproken. Er wordt een notatie gehanteerd die in de literatuur gebruikelijk is. Het is jammer dat er geen link gelegd wordt tussen regressie en variantie-analyse, want in de hoofdstukken over regressie komt men op de computeroutput vaak de term "analysis of variance" tegen. Waar de restricties vandaan komen, die opgelegd worden aan het behandelings-, en blokeffect, blijft hierdoor nu ook onduidelijk.

Het laatste hoofdstuk (H16) gaat over niet-parametrische statistiek. De bekende toetsen van Wilcoxon en Spearman's rangcorrelatietoets worden behandeld. Ik mis de verdelingsvrije variant van ANOVA (Kruskal-Wallis).

Kortom; voor iedereen die iets wil weten van de toepassingsgerichte statistiek is dit boek een goede keuze. Het bevat veel opgaven uit allerlei disciplines en geeft ook de mogelijkheid om een statistisch pakket te leren kennen. De notatie is zo gekozen, dat een aansluiting met



vervolgliteratuur niet lastig hoeft te zijn. Door de bijzonder originele opbouw is dit boek een goed vervolg op de VWO Wiskunde A stof.

F. Huele

UvA

**Morgenthaler, Stephan & Tukey, John W. (eds.).**

*Configural polysampling*

John Wiley & Sons, New York, 1991, xiv + 218 pag., ISBN 0-471-52372-0, £ 32.15.

Het lijkt me nuttig in de eerste zin van deze bespreking al te vermelden dat dit boek over robuuste schatters handelt. Om te beginnen zijn de acteurs niet tevreden met de gebruikelijke definitie van robuustheid die slechts rekening houdt met kleine afwijkingen van één centraal model. Asymptotisch (en voor grote steekproeven) mag dit een goed uitgangspunt zijn, voor kleine steekproeven achten de auteurs het ongeschikt. Zij hanteren een andere definitie, waarbij twee (sterk) verschillende modellen beide centraal staan. Probleem is dan die schatter op te sporen die voor beide modellen tegelijk optimale eigenschappen vertoont. Diverse criteria voor deze 'biopoptimaliteit' worden ontwikkeld.

Alleen symmetrische verdelingen worden in de beschouwingen betrokken, waarbij vaak het symmetriepunt de onbekende (locatie)parameter is. Daardoor is het voldoende alleen invariante schatters te bekijken; bij de efficiencyberekeningen is de Pitman-schatter het ijkpunt. Voor de twee concurrerende modellen wordt als regel de normale en de 'slash'-verdeling genomen; de laatste heeft de dichtheid

$$g(t) = \frac{1}{t^2 \sqrt{2\pi}} [1 - \exp(-t^2/2)], \quad -\infty < t < \infty.$$

Daarnaast komen het geval van een onbekende schaalparameter en de regressiesituatie aan de orde; ook robuuste betrouwbaarheidsintervallen worden behandeld.

Het boek zit vol originele ideeën en (voor mij) onbekende begrippen.

Om me te beperken tot de titel: een configuratie is een verzameling van steekproeven die - afgezien van bepaalde kengetallen - identiek zijn; voorbeelden zijn steekproeven van identieke 'scores' (m.b.v. geschat gemiddelde en standaardafwijking gestandaardiseerde steekproefwaarden). 'Polysampling' is het gebruik van steekproeven om schattingen te verkrijgen bij verschillende achterliggende modellen; dit wordt mogelijk door een geschikte weging van waarnemingen.

Voor de berekening (en schatting) van de efficiency van de diverse schatters blijkt flink wat techniek en rekenwerk nodig te zijn. De ondertitel van het boek luidt;: A route to practical

robustness; over de lengte van de route laat ik dan ook maar bier uit. Zeker is dat hier een boeiend onderzoeksgebied is aangeboord.

Hoewel in totaal zes auteurs aan het boek hebben bijgedragen, is er toch een goed samenhangend geheel ontstaan. Dit zal zeker mede komen doordat 3 auteurs samen acht (van de twaalf) hoofdstukken en de Appendix voor hun rekening hebben genomen. Naast de editors (SM en JT) is dit Katherine Bell Krystinik (KK). Een opsomming van de verschillende bijdragen volgt tot slot.

1. Background (SM & JT)
2. The background of Configural polysampling: a historical perspective (Michael Cohen)
3. Key ideas and outline (SM & JT)
4. Point estimation: location-and-scale configurations (KK & SM)
5. Compromises in the face of two situations (KK)
6. Point estimation of location: criticism and improvement (KK & SM)
7. Point estimation of location: technical choices (SM & JT)
8. Conditional confidence intervals for a location parameter (SM)
9. Confidence intervals for a scale parameter (SM)
10. Regression (Fanny O'Brien).
11. Location compromise maximum likelihood estimators (George Easton)
12. Examples (SM)

Appendix. A circular approach to cubature (JT)

J.J.A. Moors  
KUB

**Ashish Sen en Muni Srivastava**

*Regression analysis: theory, methods and applications*

Springer-Verlag, Berlijn, 1991, xv + 347 pag., ISBN 3-540-97211-0, DM 88.00.

In het voorwoord schrijven de auteurs dat dit boek een up-to-date beschrijving geeft van de theorie en de methodes van regressie analyse. Na bestudering van dit boek zal de lezer een grondige kennis bezitten van de theorie en zal hij tevens de ervaring hebben om regressie analyse in de praktijk toe te passen.

Het boek is verdeeld in twaalf hoofdstukken en drie appendices. In het eerste hoofdstuk wordt aan de hand van een voorbeeld en plaatjes uitgelegd wat regressie analyse is en wanneer en waarvoor je het kunt gebruiken. De kleinste kwadraten schatters en de Gauss-Markov condities komen aan bod. In het tweede hoofdstuk wordt het eerste hoofdstuk nog eens overgedaan voor het multiële regressie model. Hoofdstuk 3 gaat over betrouwbaarheidsintervallen en het toetsen van hypothesen. Er worden toetsen voor de coëfficiënten getoond en de



Likelihood Ratio toets. Verder worden betrouwbaarheids-intervallen voor onder andere de verwachting van de voorspelde waarde en voor de regressie parameters vermeld. Indicatoren en dummyvariabelen komen in hoofdstuk 4 aan bod. De volgende hoofdstukken gaan over het breken van de veronderstellingen van het klassieke lineaire regressiemodel: hoofdstuk 5 behandelt de aanname van normaal verdeelde residuen, hoofdstuk 6 gaat over heteroscedasticiteit en hoofdstuk 7 over gecorreleerde storingen. In deze hoofdstukken komen onderwerpen als Weighted Least Squares, Generalized Least Squares en variantie stabiliserende transformaties aan de orde. Verder geneste storingen en correlatie in de ruimte en de tijd. Hoofdstuk 8 gaat over uitbijters: wat is een outlier en hoe vind je ze. In hoofdstuk 9 worden transformaties behandeld. Gekeken wordt naar het multiplicatieve model, het logit model en polynomen. Verder worden verschillende methodes behandeld om te beslissen of, en zo ja welke transformatie gebruikt moet worden. Hierbij worden plotjes bekeken en Box-Cox transformaties. Ook wordt gekeken naar het toevoegen van termen en niet-lineariteit. Hoofdstuk 10 gaat over multicollineariteit: de effecten en het opsporen ervan. In hoofdstuk 11 komt het probleem van variabelen selectie aan bod. Er wordt ingegaan op het effect van het laten vallen van variabelen en er worden verschillende procedures beschreven om variabelen te helpen selecteren (stepwise en stagewise procedure, backward elimination en forward selection). Hoofdstuk 12 gaat over onzuivere schattingen, Principale componenten regressie, Ridge regressie en Shrinkage schatter.

In het boek worden heel veel voorbeelden behandeld die meestal gebaseerd zijn op echte data. Bij het boek wordt een diskette geleverd met een aantal data sets. Er wordt veel gewerkt met plotjes om dingen uit te leggen. Op de titel pagina staat trots vermeld dat er 38 illustraties zijn opgenomen. Aan het einde van elk hoofdstuk staat een behoorlijk aantal opgaven. Voor het maken van deze opgaven is kennis van een Statistisch software pakket nodig, in het boek wordt veel verwezen naar SAS en MINITAB. Enige kennis van statistiek en lineaire algebra is ook noodzakelijk, alhoewel de belangrijkste aspecten hiervan in de appendices worden behandeld. Het boek gaat praktisch alleen over lineaire regressie, alleen in de laatste appendix wordt een stukje niet-lineaire regressie behandeld.

Ik denk dat dit boek een redelijke inleiding in regressie analyse geeft. De grondbeginselen worden helder uitgelegd en de vele voorbeelden maken duidelijk waar de problemen kunnen optreden. Het boek is voornamelijk ingesteld op de praktijk; veelal worden problemen met behulp van procedures uit software pakketten opgelost. Ik ben het dan ook met de auteurs oneens dat een *grondige* kennis van de theorie wordt verkregen. De opgedane theoretische kennis is zeker niet voldoende voor een in de theorie van de regressie analyse geïnteresseerde persoon.

Carla van den Oudenrijn  
UvA

**Joep Brinkman***Cijfers spreken, statistiek en methodologie voor het hoger onderwijs*

Wolters-Noordhoff, Groningen, 1992, 327 pag., ISBN 90-01-16618-0, f 54,-.

Dit boek verhaalt op aardige wijze hoe onderzoek wel en hoe het niet moet worden aangepakt. Daarbij wordt wat statistiek uit de doeken gedaan, maar niet zeer diepgaand. Als leerboek statistiek is het onvolledig, vooral vanwege de geringe hoeveelheid vraagstukken. Door de vlotte stijl en de nadruk op de band tussen statistiek en onderzoek is het toch zeer aanbevelenswaard. De meeste aandacht gaat uit naar onderzoeksopzetten waarbij mensen de onderzoeksobjecten zijn, zoals bij de sociale wetenschappen en marktonderzoek. Mogelijke fouten in interpretatie en presentatie van gegevens worden in ruime mate belicht.

Voordat Brinkman cijfers laat zien of spreken, geeft hij circa 100 bladzijden tekst. Daarin worden begrippen als variabele, schaal, meten, en dergelijke uitvoerig uit de doeken gedaan. De problemen bij het vaststellen van de problemen ('probleemstellingen') en het operationaliseren van variabelen en meetmethoden worden besproken. Brinkman definieert zorgvuldig wat hij met verschillende termen bedoelt: zo maakt hij onderscheid tussen waarnemen en observeren (hoe zou men dat in de Engelse literatuur oplossen?) en tussen doelen en motieven van onderzoek. Ofschoon men het niet met de gebruikte definities eens hoeft te zijn, is een consequent gebruik van de door hem gedefinieerde woorden wel belangrijk. Sommige hoofdstukken zijn enigszins filosofisch en wijdlopij, maar naar mijn idee toch zeer geschikt om studenten de ideeën duidelijk te maken.

Elk hoofdstuk is voorzien van 'opdrachten', dat wil zeggen vragen, opdrachten of vraagstukken. Met name in hoofdstuk 3 zijn de bedoelingen van de vragen niet altijd even duidelijk. Gelukkig bevat het boek tevens antwoorden of uitwerkingen, zodat na lezing daarvan de gedachtengang van de auteur nog wel te achterhalen is. Een enkele keer verschilde ik duidelijk van mening over het antwoord op zo'n dubieuze vraag. Als van 'duurzaamheid' en 'prijs' van fietsbanden de prijs als de onafhankelijke variabele wordt aangemerkt en de student koopt een dure fietsband, omdat hij veronderstelt dat die langer meegaat, dan denkt hij vermoedelijk hetzelfde als de auteur. Als hij echter na zijn studie fietsenhandelaar wordt, en zijn banden duurder maakt in de hoop dat ze daardoor langer meegaan, dan hadden dergelijke opgaven beter niet in dit boek kunnen staan. En dit laatste vind ik van meer opgaven. De student is nog niet zover dat hij met auteurs van mening durft te verschillen, dus twijfelachtige opgaven moeten ontbreken.

De inhoud bevat:

- Het belang van onderzoek en de manier waarop de vragen tot stand komen (hoofdstuk 1 en 2)
- Variabelen en meten (hoofdstuk 3 en 4)
- Methoden voor het verwerven van gegevens (hoofdstuk 5)
- Steekproeven (hoofdstuk 6)
- Kengetallen bij variabelen, beschrijvende statistiek (hoofdstuk 7 en 8)



- Kansrekening (hoofdstuk 9)
- Enkele toetsen en betrouwbaarheidsintervallen (hoofdstuk 10)
- Interpretatie (hoofdstuk 11, 12 en 13)
- Tabellen:

Kansen uit de standaardnormale (voor positieve en negatieve  $z$ !) en binomiale verdeling (voor  $p < \frac{1}{2}$  en voor  $p > \frac{1}{2}$ !), kritieke waarden uit de chi-kwadraat-verdeling, de  $t$ -verdeling en de verdeling van de steekproefcorrelatiecoëfficiënt, de betrouwbaarheidsvis voor de binomiale betrouwbaarheidsintervallen.

Een opsomming van de met enige moeite vergaarde minpunten:

- Sommige vragen zijn niet erg duidelijk (maar in combinatie met de gegeven antwoorden wel leuk en leerzaam).
- De hoeveelheid statistiekvraagstukken is beperkt (één opgave moet voldoende zijn om de student met de berekening van kansen uit de normale verdeling vertrouwd te maken).
- De tekst bevat enkele hinderlijke germanismen en andere rare woorden (overstijgen i.p.v. te boven gaan, uitzuiveren, inschatten i.p.v. schatten, aanname i.p.v. veronderstelling, belvormig i.p.v. klokvormig, meekrijgen i.p.v. begrijpen, binomiaalverdeling i.p.v. binomiale verdeling). Er zijn ook enkele detail-fouten van statistische aard.
- Letters hier en daar wat flets gedrukt.

Veel interessanter is de lijst met pluspunten:

- De verteltrant is boeiend en bevat leuke, historische voorbeelden.
- Voetangels en klemmen komen goed aan de orde.
- De benadering van de statistiek geschiedt vanuit de praktijksituatie, getallen en berekeningen komen op de tweede plaats. (Omdat de statistiek niet zeer uitvoerig wordt behandeld zou de ondertitel beter hebben kunnen luiden: Methodologie met statistiek ... enz.)

Al met al: als onderdeel van een studie waarin veel statistiek voorkomt, zeker aanbevolen literatuur.

J.M. Buhrman

Hogeschool voor Economische Studies, Amsterdam

**Prem S. Mann**

*Introductory Statistics*

John Wiley & Sons, New York, 1991, xix + 774 pag., ISBN 0-471-52733-5, £ 45.95.

Dit leerboek voor eerste lessen in toegepaste statistiek is bedoeld voor studenten zonder sterke wiskundige achtergrond. Het boek maakt deel uit van een complete leermethode, waar o.a. een

instructor's manual bij hoort. Na het doorwerken van de theorie wordt de lezer vertrouwd gemaakt met het gebruik van het statistische programma MINITAB. Er zijn 13 hoofdstukken en elk hoofdstuk bevat:

- Case studies
- Opgaven per paragraaf en over het hele hoofdstuk
- Samenvatting van alle begrippen en belangrijke formules
- Hoe MINITAB te gebruiken bij dit hoofdstuk.

De opbouw is als volgt:

#### *H1: Introduction*

Introductie van alle basisbegrippen zoals onder andere (doel)populatie en (aselecte) steekproef.

#### *H2: Organizing data*

Hier worden (grafische) technieken gegeven voor kwalitatieve data; frequentieverdeling, staaf- en taartdiagram. En ook voor kwantitatieve data; frequentieverdeling in klassen, histogram, polygon en 'stem-and-leaf displays'.

#### *H3: Numerical descriptive measures*

Invoeren van het populatie gemiddelde  $\mu$ , het steekproef gemiddelde, de mediaan, de mode, range, standaard deviatie en variantie. Ook de 'Box-and-Whisker plot' komt aan de orde.

#### *H4: Probability*

Defenities van experiment, uitkomst en gebeurtenis ( $E_i$ ). Introductie van  $P(E_i)$ , voorwaardelijke kansen en kansbomen.

#### *H5: Discrete random variables and their probability distributions*

Binomiaal en Poisson verdelingen worden gegeven en het gebruik van tabellen van deze verdelingen wordt behandeld. (tabellen staan achter in het boek)

#### *H6: Continuous random variables and the normal distribution*

#### *H7: Sampling distributions*

#### *H8: Estimation of the mean and proportion*

Er worden hier twee soorten steekproeven gedefinieerd: grote steekproeven (steekproefgrootte  $> 30$ ) en kleine steekproeven (steekproefgrootte  $< 30$ ). Voor grote steekproeven wordt de normale benadering gebruikt en voor kleine steekproeven de t-verdeling. Ook komen  $100(1-\alpha)\%$  betrouwbaarheidsintervallen aan bod.

#### *H9: Hypothesis tests about mean and proportion*

Dit hoofdstuk gaat over het testen van een nul hypothese tegen een alternatieve hypothese.

#### *H10: Estimation and hypothesis testing: two populations*

#### *H11: Chi-square tests*

Voor gebruik bij kruistabellen en de populatie variantie.

#### *H12: Simple linear regression*

#### *H13: Analysis of Variance*

De F-verdeling komt hier aan de orde en ANOVA tabellen worden behandeld.



Het boek bevat zeer veel uitleg die wordt ondersteund door duidelijke plaatjes en voorbeelden, zodat de stof makkelijk te behappen is. Resultaten, definities en stellingen zijn in aparte kaders duidelijk te traceren, maar de afleidingen en bewijzen ontbreken in het boek.

In H1 wordt de sommatie-notatie als volgt ingevoerd:  $\Sigma x = x_1 + \dots + x_n$ . Het weglaten van de indices bij het sommatie teken leidt later tot verwarring, als in één kader de formules voor het populatie- en steekproefgemiddelde als volgt worden weergegeven:

$$\mu = \frac{1}{N} \sum x \quad \text{en} \quad \bar{x} = \frac{1}{n} \sum x$$

Er worden hier natuurlijk twee verschillende  $\Sigma x$ -en mee bedoeld. In het algemeen had de notatie in het boek wel wat nauwkeuriger mogen zijn. Ook de kansdichtheden van de normale-, t-, F- en  $\chi^2$  verdeling ontbreken. Tussen stochasten en de waarden hiervan wordt geen onderscheid gemaakt. Het boek is dan ook te simpel opgezet voor mensen die op een wiskundige manier naar de stof willen kijken.

Al met al is dit boek volgens mij zeer geschikt voor de doelpopulatie. De grote hoeveelheid voorbeelden, opgaven en de goede layout maken dat dit boek makkelijk weg leest en begrijpbaar blijft. Ook het integreren van het gebruik van een statistisch programma (MINITAB) lijkt mij zeer nuttig in een cursus toegepaste statistiek.

Sylvia de Vries

UvA

### **T.J. Hastie & R.J. Tibshirani**

#### *Generalized Additive Models*

Chapman and Hall, London, 1990, xv + 335 pag., ISBN 0-412-34390-8, £ 25.00.

Dit is een belangrijk en boeiend boek waarin een zeer toepasbare techniek beschreven wordt. Helaas is het boek niet helemaal toegankelijk voor de gemiddelde gebruiker, die zal op een eenvoudiger boek of computerpakket moet wachten. Voor degene die zich thuisvoelt in de statistiek en de kansrekening is het boek aan te bevelen.

Gegeneraliseerde additieve modellen zijn vanuit twee goed bekende ideeën ontwikkeld, namelijk, *regressie analyse* en *smoothing*. In principe trekken de modellen een kromme of oppervlakte door een puntenwolk. Trends en patronen worden weergegeven door grafieken. Daardoor zijn de doelen en principes makkelijk te begrijpen, in tegenstelling tot de algoritmen en eigenschappen die moeilijker zijn. Bovendien zijn de technieken niet echt bruikbaar zonder computergrafieken. Desalniettemin, als een ondersteunend pakket beschikbaar is, bieden de technieken veel mogelijkheden aan om data (van zelf) voor zich te laten spreken, data te

onderzoeken, en patronen op te sporen. Het uitgangspunt is de vraag *hoe kan de houding tussen verklarende variabelen en afhankelijke variabelen als een model beschreven worden*. De auteurs nemen aan dat de gegevens de geschikte vorm zouden kunnen aantonen en daarvoor gebruiken ze de slagzin: *let the data show us the appropriate functional form*.

De ontwikkeling is als volgt. Het lineaire model voor een één-variabele regressie,

$$E(Y|X) = \alpha + \beta x$$

wordt vervangen door

$$E(Y|X) = f(X)$$

multi-variate regressie door additieve modellen

$$E(Y|X) = \alpha + \sum_{j=1}^{j=p} f_j(x_j)$$

Ten laatste wordt de *generalized additive model* als

$$\mu = E(Y|X)$$

$$g(\mu) = \alpha + \sum_{j=1}^{j=p} f_j(x_j)$$

gegeven. De tussenstap is in het boek *Generalized Linear Models*<sup>1</sup> te vinden. De functie  $g$  is *link functie* genoemd en de  $f_j$  *smoothers* omdat die gegeneraliseerde versies van moving averages zijn. Dus worden de modellen algemener dan lineaire regressie behalve dat interacties buiten beschouwing blijven. Het belangrijkste is dat de modellen niet tot lineaire vorm zijn beperkt.

Het verhaal is te vinden in de hoofdstukken 2, 4 en 6 en technische details in de andere hoofdstukken. In hoofdstuk 2 wordt het idee van een "smoother" in het één-variabele geval uitgelegd om de latere ontwikkeling te motiveren. In hoofdstuk 4 wordt het additieve model geïntroduceerd en in hoofdstuk 6 worden de gegeneraliseerde additieve modellen geïntroduceerd. Een aantal bekende technieken zijn terug te vinden als bijzondere gevallen van gegeneraliseerde additieve modellen, waaronder logistieke regressie, Box-Cox transformaties, proportional hazards, proportional odds. In de resterende hoofdstukken worden de technische details gehanteerd, die diepgaande kennis eisen van het gebruik van waarschijnlijkheidsfuncties, splines, Hilbertruimte, en lineaire algebra. In het laatste hoofdstuk worden twee cases uitgebreid beschreven.

De mogelijkheden binnen gegeneraliseerde additieve modellen, in het bijzonder de grafieken, maken de lezer enthousiast om de modellen toe te passen, maar de technische eisen maken dit



moeilijk. Helaas zijn er momenteel geen computer pakketten te krijgen, behalve een van de auteurs, en wat functies in de S taal.

#### Referenties

- [1] McCullagh, P. and Nelder, J.A. (1989). *Generalized Linear Models*. Chapman and Hall, Londen.

Martin Newby

Faculteit der Bedrijfskunde, TUE

#### **T.H.B. Hendriks en P. van Beek**

##### *Optimaliseringstechnieken: principes en toepassingen (3e druk)*

Bohn Stafleu Van Loghum, Houten, 1991, ix + 348 pag., ISBN 90-313-1241-X, f 49.00.

The book is now available in its 3rd, completely revised edition (1st edition: 1983). It addresses the topics linear programming (LP), transport problems (TP), integer programming (ILP), dynamic programming (DP), and non-linear programming (NLP). The reader is thus introduced into most of the basic optimization techniques of operations research (OR). This is a well thought out choice of material. But I have some questions with regard to its presentation, which appears somewhat unbalanced so that neither the academic researcher nor the non-academic user might feel being fully taken care of.

An introductory chapter discusses OR applications. Unfortunately, many of the modelling problems mentioned are worked out (e.g., it is stated that Example 2.7 "can be formulated" as a mixed integer LP -an explicit formulation, however, is not provided). The travelling salesman problem (TSP) is given as an example for ILP. The chapter on ILP, however, only lists the general methods of Gomory cuts and branch and bound (with LP-relaxation) as solution techniques, which are not the best ways to approach the TSP with its particular structure in practice. Problems giving rise to NLP-models are missing completely.

The emphasis of the chapter on LP is on the theory of the simplex method (the ellipsoid method, Karmakers's method and the use of LP software are briefly sketched). After a review of linear algebra, the basic properties of LP are theoretically derived and lead to stepwise accumulated information on the simplex method until finally the simplex method can be stated. This approach is exactly opposite to the one in, e.g., Chvatal ("Linear Programming", Freeman & Co., 1983), where the simplex method is introduced right away and then used as a tool to derive the basic properties of LP algorithmically.

The simplex algorithm for transport problems in the following chapter is developed within the framework of LP without graphtheoretic interpretations. Generally, graphs do not play a major role in the book and combinatorial optimization as such is not emphasized. As already mentioned above, the two general techniques presented for dealing with integer linear programs

are a bit restrictive in practice. Here one could wish that more powerful methods (Lagrange relaxation, heuristics) were also included.

I like the inclusion of dynamic programming (both deterministic and stochastic) into the book. No other Dutch introduction into OR that I am aware of offers a similar treatment of this important topic. The final chapter is devoted to NLP. The relevant notions and facts from analysis are stated and subsequently the basic techniques of NLP are described. The reader who just wants to get a quick overview of the techniques is in this chapter served best because the authors abstain from packing in too much theory.

In summary, I think that we have here a useful book that introduces the reader to the basic techniques of OR. The style is somewhat dry in that facts are listed more or less one after the other. (At least this reviewer would not mind more internal motivation behind the presentation). The reader of the book will gratefully appreciate the care with which proof reading seems to have been done (I did not detect a single error). Finally, it should be mentioned that another Dutch book (Dirickx, Baas, Dorhout: "Operationele Research", Academic Service, 1987) is devoted to about the same topics (combinatorial optimization is more stressed while dynamic programming is not discussed). The interested reader may find there additional examples and problems worked out.

U. Faigle

Faculteit der toegepaste wiskunde TUT

**Geoffrey J. McLachlan**

*Discriminant analysis and statistical pattern recognition*

John Wiley & Sons, New York, 1992, xv + 526 pag., ISBN 0-471-61531-5, £ 52.00.

Dit lijvige boekwerk (446 bladzijden tekst, 60 bladzijden bibliografie en 20 bladzijden indices) beoogt een overzicht te geven van de stand van zaken op het gebied van discriminantanalyse anno 1990. Zoals de auteur in zijn voorwoord betoogt, is het niet zijn bedoeling om een leerboek te schrijven. Dat is ook aan de stijl van het boek te merken.

Compact wordt alles beschreven in een vrij formele wiskundige stijl. Voorbeelden zijn er wel, maar niet overvloedig en kritische beschouwingen ontbreken meestal. Men zou het boek kunnen karakteriseren als een uitstalkast waarin alles staat wat op dit terrein te koop is en waaruit de lezer zijn keuze kan maken. De goede indices naar auteur en trefwoorden zijn daarbij zeer behulpzaam. De auteur heeft gestreefd naar volledigheid, wat misschien nog het duidelijkst wordt in de auteurindex, waaruit af te lezen valt dat het merendeel van de uitgebreide bibliografie, misschien zelfs wel alles, in de tekst ook geciteerd wordt. Dat maakt de tekst erg vol en heeft tot gevolg dat ondanks de lijvigheid elk onderwerp op zich zeer beknopt



behandeld wordt. Een plezierige consequentie van de volledigheid is dat werk van landgenoten uitgebreid genoemd wordt. Met name het werk van Habbema & Hermans, Schaafsma & Ambergen en Schmitz krijgt ruime aandacht, terwijl van onze zuiderburen ook Albert & Lesaffre ruime aandacht krijgen.

Ook mijn artikel met S. le Cessie uit *Statistica Neerlandica* 1988 wordt geciteerd. Bewijs ten overvloede dat SN een tijdschrift is dat ook buiten Nederland gekend wordt.

Na deze poging mijnerzijds de stijl van het boek te beschrijven, een korte beschrijving van de inhoud.

De eerste 6 hoofdstukken geven een extensieve en gedetailleerde beschrijving van de klassieke discriminant-analyse gebaseerd op multivariate normale verdelingen met alle wiskundige finesses die daar bij horen. De eerste voorbeelden verschijnen op pagina 201. Het allereerste voorbeeld is de vaststelling van draagsterschap van Haemofilie A, gebaseerd op Hermans & Habbema (1975).

Hoofdstuk 7 bespreekt discrete en gemengde data, nog steeds vanuit de discriminantanalyse optiek. In hoofdstuk 8 komt de logistische regressie aan de orde, maar dan vooral als (niet altijd efficiënt) alternatief voor discriminantanalyse. In hoofdstuk 9 worden niet parametrische methoden als kernel-methodes (ALLO) en "classification-trees" (CART) besproken. De performance van de gepresenteerde classificatieregels komen aan de orde in hoofdstuk 10, waar het schatten van de "error rate" via cross-validatie en bootstrap wordt besproken. Pas in hoofdstuk 11 komt de onnauwkeurigheid van de geschatte a posteriori kans aan de orde. Hoofdstuk 12 bespreekt het effect van de selectie van variabelen.

Daarmee is het gehele scala van de discriminantanalyse aan de orde geweest. Dat in hoofdstuk 13 nog "Statistical Image Analysis" aan de orde komt is begrijpelijk vanuit het oogpunt van public relations voor de statistiek, maar het is toch een beetje vreemde eend in de bijt.

Resumerend: Een gedegen boek, met een schat aan informatie voor de beoefenaars van de statistiek. Het geeft een zeker gevoel van trots over de produktiviteit en inventiviteit van de professie, maar ik mis iets van het gevoel voor de schoonheid van het vak dat boeken die wel leerboeken durven zijn, bij mij soms te weeg kunnen brengen.

J.C. van Houwelingen

Faculteit de Geneeskunde, RUL

Carl-Erik Särndal, Bengt Swensson, Jan Wretman

*Model Assisted Survey Sampling/Springer Series in Statistics*

Springer-Verlag, Berlijn, 1992, xv + 694 pag., ISBN 3-540-97528-4, DM 98.00.

Model Assisted Survey Sampling is een belangwekkend boek. Hoewel stijl en onderwerp afwijken van het in 1977 verschenen boek *Foundations of Inference in Survey Sampling* van Cassel, Särndal en Wretman, is de grondigheid waarmee het boek is geschreven vergelijkbaar. Het nieuwe boek is vooral ook een boek voor de praktijk. Hierbij ligt de nadruk op algemeen bruikbare methoden voor het schatten van populatieparameters. Uitgangspunt is steeds de steekproef waarbij de elementen met gegeven insluitkansen zijn getrokken. Hulpinformatie bekend voor de gehele populatie wordt met behulp van regressieachtige technieken gebruikt om tot goede schatters te komen. Bekende schatters als de quotiënt- en de regressieschatter worden op een natuurlijke wijze afgeleid uit algemene principes. Dit gebeurt door het hele boek heen. Hierdoor ontstaat er een grote samenhang tussen de verschillende onderdelen van het boek. Het boek bestaat uit vier delen. De hoofdstukken in het boek eindigen steeds met een uitgebreide verzameling opgaven, soms met antwoorden.

Deel I gaat in op de grondbeginselen van het schatten van parameters van eindige populaties en behandelt de meest gebruikte enkel- en meervoudige steekproefdesigns. In vijf hoofdstukken worden theorie en praktijk ingeleid, en wordt een basis gelegd voor schatters bij kanssteekproeven, waarbij de Horvitz-Thompson schatter een grote rol speelt. In hoofdstuk 5 wordt de Taylorreeksbenadering voor het schatten van varianties op een heel duidelijke manier ingeleid.

Deel II legt de basis voor het onderwerp van het boek: schatten door lineair modelleren gebruik makend van hulpvariabelen. Er wordt diepgaand ingegaan op de regressieschatter, waarvan de algemene gedaante op maar liefst zes verschillende wijzen elk met een eigen interpretatie is te schrijven. Aan dit onderwerp worden drie hoofdstukken gewijd, waarin ook toepassingen bij verschillende steekproefdesigns worden behandeld.

Deel III behandelt in vijf hoofdstukken onderwerpen als de twee-fasen steekproef, het schatten voor deelpopulaties, het schatten van varianties met behulp van methoden als de balanced half-sample techniek, het ontwerpen van optimale designs en analyse van steekproefgegevens.

Deel IV besluit het boek in vier hoofdstukken die gaan over niet- steekproeffouten, nonrespons, meetfouten en rapportage over de kwaliteit van steekproefuitkomsten.



Het boek is ten eerste aan te bevelen voor ieder die zich serieus bezig houdt met de theorie en praktijk van steekproeven. De inspanning die het kost de bijna 700 dichtbedrukte bladzijden te bestuderen wordt beloond.

Jos de Ree

CBS

**A.S. Hedayat & Bikas K. Sinha**

*Design and inference in finite population sampling*

John Wiley & Sons, New York, 1991, xiii + 377 pag., ISBN 0-471-88073-6, £ 39.95.

De auteurs hebben het boek opgebouwd uit lesmateriaal dat zij hebben ontwikkeld voor hun colleges aan vele universiteiten. Zij noemen het een inleiding in de theorie van het steekproefonderzoek. In eerste instantie hebben zij als lezers universitaire studenten voor ogen, maar ook serieuze gebruikers van steekproefmethoden worden geacht profijt te kunnen hebben van de theoretische resultaten. Zij veronderstellen bij hun lezers waardering voor subtiele wiskundige en statistische redeneringen.

Het boek bevat twaalf hoofdstukken, elk met oefeningen en een literatuur lijst. In 1 en 2 worden de basisbegrippen behandeld, alsmede enige existentiële stellingen. Terloops komt daarin ook de "celebrated" Horvitz-Thompson-schatter (HTE) aan de orde. Hoofdstuk 3 is geheel gewijd aan de eigenschappen van de HTE, het schatten van varianties en met name die van de HTE. In 4 wordt simple random sampling (SRS) en enige varianten daarvan besproken. Hoe hulpinformatie in designs en schatters kan worden verwerkt is in 5 en 6 beschreven. In 7, 8 en 9 worden achtereenvolgens cluster-, systematische en gestratificeerde steekproeven behandeld. In 10 en 11 komen superpopulaties en randomized respons-technieken bij delicate vragen aan de orde. Tenslotte worden in 12 nog de onderwerpen non-respons, small area estimation en resamplingtechnieken kort aangestipt.

Direct valt het fraaie zetwerk op. Dit maakt het bestuderen van het boek tot een groot genoegen. De theorie wordt gebracht op hoog niveau, waarbij in de literatuur nog onopgeloste problemen ook worden vermeld. In afwijking van de traditionele aanpak worden uitsluitend steekproeven beschouwd waarin de elementen hoogstens eenmaal voorkomen. Toch wordt SRS met teruglegging wel behandeld, maar op een nogal gekunstelde wijze als mixture design. De situaties waarin een bepaalde theorie toepasbaar is worden helder aangegeven, bijvoorbeeld wanneer superpopulaties een rol kunnen spelen. De tekst bevat bovendien vele duidelijke voorbeelden en overzichten. Het is een zeer waardevol boek over een onderwerp waarvan de

auteurs zelf in hun inleiding zeggen: "The subject of survey sampling holds rich rewards for those who learn its secrets."

Elly Koeijers

Hoofdafdeling Statistische Methoden CBS, Voorburg.

### **J. Muilwijk, T.A.B. Snijders en J.J.A. Moors**

#### ***Kanssteekproeven***

Stenfert Kroese/Martinus Nijhoff, Leiden, 1992, xv + 282 pag., ISBN 90-207-2221-2, f 59.90

Het hier te bespreken boek is een gemoderniseerde en uitgebreide versie van "Steekproeven, een inleiding tot de praktijk" van Moors en Muilwijk (1975, Agon Elsevier). Evenals zijn voorganger is het een voor Nederlandstalige statistici, methodologen en onderzoekers gemakkelijk toegankelijke inleiding tot de theorie van steekproefonderzoek. Hiermee biedt het boek soelaas voor beginnende wetenschappers die niet erg goed uit de voeten kunnen (een dikwijls miskend probleem!) met soortgelijke Engelstalige boeken, zoals "Sampling Techniques" van Cochran (1977, Wiley). Wie behoefte heeft aan verdere verdieping wordt de weg gewezen naar de dikwijls Engelstalige literatuur. Het zou in dit verband gemakkelijk geweest zijn als de Engelstalige equivalenten van enkele kernbegrippen in een apart lijstje waren opgenomen in plaats van in de tekst.

Het boek is strakker en wiskundiger van stijl dan Moors en Muilwijk (1975). Het is hiermee niet minder toegankelijk geworden, want de theorie wordt zonder uitvoerige bewijzen behandeld. De geïnteresseerde lezer kan aan het einde van een hoofdstuk belangrijke wiskundige uitwerkingen vinden in een aparte paragraaf, voorafgegaan door aanvullende opgaven en gevolgd door antwoorden. Het boek bestaat uit de delen A De bekendste steekproef technieken

B Algemene theorie en diverse technieken

C Steekproeven in de praktijk

Deel C geeft een overzicht van de behandeling van niet-steekproeffouten en een globale 'procesbeschrijving' van steekproefonderzoek. Voor uitvoerige praktische voorbeelden wordt de lezer echter naar de literatuur verwezen.

Bij elke steekproefstrategie (ontwerp en schattingsmethode) wordt uitvoerig ingegaan op allerlei afwegingen die relevant zijn voor de keuze van deze strategie, ondermeer afhankelijk van de aard van eventueel beschikbare hulpinformatie. Dit gebeurt overzichtelijker dan b.v. bij Cochran (1977) het geval is. In deel A komen zo in een logische volgorde enkelvoudige, gestratificeerde en clustersteekproeven aan de orde, alsmede systematische en getrapte



steekproeven. Trekking met ongelijke kansen ('evenredig aan de omvang') met teruglegging wordt hierbij als alternatief gepresenteerd voor gestratificeerde steekproeven met ongelijke kansen *zonder teruglegging*. In een logische verstrengeling met deze opbouw worden de bekendste schattingsmethoden geïntroduceerd: directe schatter, verschillschatter, quotiënt schatter en regressieschatter. Trekking met ongelijke kansen zonder teruglegging wordt in deel B behandeld. Jammer genoeg wordt hierbij niet gewezen op de overeenkomst van de hierbij geïntroduceerde "Horvitz-Thompson schatter" met de schatter voor trekking met teruglegging uit deel A. Een in de praktijk nuttige techniek als stratificatie achteraf (post-stratificatie) komt helaas summier en versnipperd aan bod: (i) er wordt weliswaar verwezen naar de methode van 'lineair wegen' maar niet naar iteratief proportioneel fitten, (ii) poststratificatie wordt in deel A alleen gemotiveerd vanwege variantie-reductie terwijl in deel C alleen gewezen wordt op de mogelijke reductie van vertekening ten gevolge van selectieve nonrespons (het motief van 'standaardisatie' wordt niet genoemd).

Deel B bevat goede introducties tot een aantal diepgaander onderwerpen: naast trekking met ongelijke kansen zonder teruglegging ook matrixsteekproeven, variantieschatten (random group, jackknife en linearisatie) en toetsingstheorie. Het is voor de praktijk van belang dat matrixsteekproeven, in het bijzonder steekproeven in ruimte en tijd, hierbij veel uitvoeriger aan bod komen dan in vergelijkbare boeken gebruikelijk is.

Dit mooie boek is al met al een solide en heldere introductie tot een breed spectrum van steekproeftechnieken en daarbij van pas komende statistische methoden. Voor de 'gevorderde' kan het ook als efficiënt naslagwerk dienen. Enkele typfouten in formules, de verwijzing naar een ontbrekende lijst van symbolen achter in het boek, en het modernistisch ogende gebruik van afkortingen zonder punt (of acronyemen met spaties: "men trekt e a z t trossen") doen hier niet aan af.

Ger Slootbeek en Hans Akkerboom,  
Afdeling Statistische Methoden CBS, Heerlen.

### A.I. Dale

*A history of inverse probability, from T. Bayes to K. Pearson*

Springer-Verlag, New York, 1991, xx + 495 pp, ISBN 0-387-97620-5, DM 124.00.

Andrew I. Dale (Department of Mathematical Statistics, University of Natal, South Africa) schrijft in het voorwoord dat het doel van dit boek het presenteren van een overzicht over het verschenen werk betreffende 'inverse probability' is, en dat hij zich hierbij beperkt heeft tot de

periode liggend tussen de bijdragen van Thomas Bayes en Karl Pearson vanwege de omvang van het werk. De probleemstelling die met de term *inverse probability* aangeduid wordt, het bepalen van kansen van mogelijke oorzaken gegeven een aantal waarnemingen, is inderdaad voor het eerst verwoord door Bayes in zijn beroemde 'Essay towards solving a problem in the doctrine of chances', dat waarschijnlijk door slechts weinige statistici gelezen is. Het beëindigen van het boek bij Pearson is minder voor de hand liggend, maar begrijpelijk gezien de enorme toename in het aantal werken die betrekking hebben op dit onderwerp in de periode na Pearson, en de faam van deze als statisticus. Laten we maar aannemen dat het feit dat het "tijdperk Fisher" grotendeels buiten de gekozen periode valt niet van invloed is geweest.

Om de inhoud van het boek goed weer te geven hadden in de sub-titel ook de namen van Condorcet en Laplace verwerkt mogen worden, blijkens de indeling van het boek: 51 blz. zijn gewijd aan Bayes, 44 aan Condorcet, 112 aan Laplace en 15 aan Pearson. Verder bevat het boek een zeer nuttig overzicht van nagenoeg het totale werk op dit gebied dat verscheen in de genoemde periode, hetgeen veel onbekend en interessant werk toegankelijk maakt. Dit laatste feit, alsmede het nauwkeurig bij elkaar brengen van ideeën en kritieken hierop maken het boek waardevol. Helaas is het moeilijk leesbaar door de gekozen documenten-stijl, waarbij zoveel mogelijk oorspronkelijke notaties, definities en talen gebruikt zijn. Met name dit laatste is een hindernis bij het lezen. Het is aardig dat er ook een Nederlandse zin in het boek opgenomen is (in de notes), zijnde de omschrijving bij Thomas Bayes uit de Winkler Prins Encyclopedie van 1948: "Engels wiskundige omtrent wiens leven wij vrijwel niets weten".

Als Dale oorspronkelijke notaties verandert gaat het wel eens mis, zoals op blz. 220 (Laplace), waar foutieve definitie van een gebeurtenis  $E$  leidt tot  $\Pr(E \mid \text{not } E) = 1/999$ . Alhoewel de aandachtige lezer dergelijke foutjes snel ontdekt zijn ze behoorlijk irritant.

Het boek is rijk aan 'notes', die nagenoeg alle overgeslagen kunnen worden maar zeker nuttige achtergrond informatie bevatten. Vertalingen van belangrijke passages die in de tekst in oorspronkelijke taal weergegeven zijn komen hier helaas maar zelden voor. Bijzonder aardig is de appendix met een niet eerder gepubliceerde 'reading note' van Savage, geschreven als persoonlijke samenvatting van Bayes' essay, voorzien van commentaar. Het is ook leuk te lezen van welke aard de problemen waren die vroeger aan statistici voorgelegd werden.

Hoofdstukken 1, 2 en 3 beschrijven respectievelijk Bayes' leven, het essay en commentaar op het essay. Dale's eerdere wetenschappelijke publicaties hadden ongeveer alle betrekking op Bayes naar wie hij veel onderzoek verricht heeft, maar het blijft duidelijk dat er maar heel weinig over Bayes' leven bekend is. Een verdienste van Dale is het benadrukken van de rol van Price, een vriend van Bayes die na diens dood het essay heeft laten publiceren, en daarbij een gedeelte toegevoegd heeft waarin de aandacht voor kansen van oorzaken vervangen wordt door de vraagstelling wat kansen van toekomstige gebeurtenissen zijn, gegeven waarnemingen uit het



verleden. Mede door deze toevoeging worden Bayes' gedachten beschouwd als de fundamenteën van de hedendaagse Bayesiaanse statistiek. Het is zeer spijtig dat Bayes' essay zelf niet opgenomen is, terwijl het eigenlijk noodzakelijk bestudeerd zou moeten worden alvorens aan deze hoofdstukken te beginnen. Het is de geïnteresseerde lezer dan ook aan te raden eerst het boek van Press {'Bayesian Statistics: Principles, Models, and Applications', 1989} ter hand te nemen, waarin als appendix zowel het essay als een biografie van Bayes te lezen zijn. Dit laatste is prettiger te lezen dan hoofdstuk 1 van Dale's boek.

Hoofdstuk 4 beschrijft de periode 1761-1822, waarbij weer 15 blz. gewijd aan Bayes (onder andere een bespreking van een later gevonden 'notebook' met het bewijs van een 'rule' uit het essay) en Price die gewoon bij de hoofdstukken 2 en 3 thuishoren. Per auteur wordt een idee van diens werk gegeven. Voor bestudering van dat werk is het echter te beknopt (vaak blijven beredeneringen en bewijzen achterwege). Erg aardig is het vermelden van kritiek van andere auteurs op bepaalde werken, waarbij ook Dale zo nu en dan eigen kritiek toevoegd. Met regelmaat verwijst Dale, waar nodig corrigerend, naar Todhunter {'A History of the Mathematical Theory of Probability: from the time of Pascal to that of Laplace', 1865}. In de conclusies bij dit hoofdstuk wordt een argument gegeven om niet Bayes maar Lagrange te beschouwen als de 'father of inverse probability'.

Met het wijden van afzonderlijke hoofdstukken (5 en 6) aan Condorcet en Laplace wordt het belang van beider werk terecht benadrukt, alhoewel dit wellicht ook met minder bladzijden bereikt had kunnen worden. De belangrijkste werken van beiden, 'Essai sur l'application de l'analyse à la probabilité des décisions rendues à la pluralité des voix' (Condorcet) en 'Théorie analytique des probabilités' (Laplace), worden zeer uitgebreid besproken. Dale verdient een compliment voor zijn onderzoek gepresenteerd in een appendix bij hoofdstuk 6, waarin het werk van Bayes (cq. Price) en Laplace vergeleken wordt.

Laplace's 'rule of succession' wordt tot op de dag van vandaag door Bayesianen beschouwd als het correcte antwoord van het probleem: "Stel  $n$  experimenten hebben  $s$  successen en  $n-s$  mislukkingen opgeleverd. Wat is de kans op succes bij een volgend experiment (experimenten identiek en onafhankelijk)?" Laplace's antwoord,  $(s+1)/(n+2)$ , is gebaseerd op Bayes' theorie met uniforme prior (merk op dat voor een frequentistische benadering van dit probleem een aanname over het stopcriterium van het experiment nodig is). Aandacht voor een dergelijke predictieve vraagstelling troffen we ook aan bij Price's toevoegingen aan Bayes' essay, en we mogen dit probleem als één van de meest essentiële uit de statistiek beschouwen.

In hoofdstukken 7 ('Poisson to Venn') en 8 ('Laurent to Pearson') komen weer veel auteurs aan bod, waarbij het spijtig is dat Boole's ideeën over 'lower and upper probabilities' slechts in de notes vermeld worden. Interessant zijn bijdragen waarin voor het eerst begrippen als "coherence" en "subjective probability" in de literatuur verschijnen. Doordat er echter tussen vele

bijdragen slechts kleine verschillen zijn (vaak alleen qua omschrijving van een probleem) krijgt de lezer snel een gevoel van herhaling. Ook plaatst Dale zo nu en dan kleine opmerkingen die onbegrijpelijk zijn zonder het betreffende werk te kennen.

In hoofdstuk 8 staat Laplace's 'rule of succession' centraal, die leidt tot verschillende paradoxen. Een voorbeeld, naar Poisson, is: "Stel dat er twee speelkaarten (zwart of rood, populatie waaruit deze kaarten komen is onbekend) omgekeerd op tafel liggen. Eén wordt omgekeerd en blijkt zwart te zijn. Wat is de kans dat de tweede zwart is?" Met Laplace's rule wordt het antwoord  $2/3$ , gebaseerd op de prior kans  $\Pr(\text{"kaart is zwart"}) = 1/2$ . Zo redenerend volgt enerzijds, via conditionering naar de kleur van de eerste kaart, dat we a priori weten dat  $\Pr(\text{"beide kaarten zelfde kleur"}) = 2/3$ , terwijl zonder deze conditionering deze prior kans gelijk is aan  $1/2$ .

Alle hier vermelde paradoxen zouden opgelost zijn als het antwoord uit de 'rule of succession' s/n zou zijn, hetgeen in de Bayesiaanse theorie bereikt wordt door de prior te baseren op de gegevens, in plaats van een prior aan te nemen en de gegevens te gebruiken voor het vinden van een posterior. Alhoewel dit een volstrekt logische redenering is (immers, pas bij extra data gaat de statisticus nu over tot 'updating'), blijken Bayesianen (en Dale) dit niet te accepteren, met name vanwege de gewaagde conclusies die volgen bij  $s=0$  (kans op succes is 0) of  $s=n$  (kans is 1). Zinnvoller lijkt het mij voor deze gevallen toe te geven dat je extra subjectieve informatie wilt gebruiken, en voor alle overige gevallen de gegeven data te gebruiken om de prior op te baseren. Helaas heeft Dale hier niet de kans gegrepen een eigen beschouwing toe te voegen.

Laten we tot slot opmerken dat het onderwerp van dit boek, met name de filosofie hierachter, in feite door iedere statisticus persoonlijk bestudeerd zou moeten worden. Dit boek is hiertoe zeer geschikt, maar gezien de tijd die nodig is om het te lezen en over de vragen die het oproept na te denken, is het helaas ongeschikt voor (universitaire) cursussen. Waarschijnlijk zal het als geheel slechts door onderzoekers op dit gebied (en hobbyisten) gewaardeerd worden. Als bron van referenties en originele literatuur heeft het boek grote waarde.

F.P.A. Coolen

Faculteit Wiskunde & Informatica, TUE



**M.J. Crowder, A.C. Kimber, R.L. Smith and T.J. Sweeting.**

*Statistical analysis of reliability data.*

Chapman & Hall, London, 1991, xii + 250 pag., ISBN 0-412-30560-7, £ 25.00.

The goal of this book is to cover both the statistical modelling and the statistical aspects of the reliability, i.e. to fit and judge complex models to observed reliability data by numerical techniques. The book is aimed for engineers with at least a basic knowledge of elementary statistics and experiences of some statistical computer packages. The authors have tried to make the book self-contained. However, within the limited size (250 pages) of the book, the authors have tried to cover too much of advanced statistical theories, making it difficult for the practical working engineer or even for the industrial statistician to choose among the many possible methods. The frequently encountered references in the text and the given examples, make it also necessary to have access to a large amount of books and articles in statistics. It may be frustrating for the practitioner, in the industrial world, not having copies on hand of the many references when reading the book. Moreover, some of the chapters contains theories and procedures not normally available in statistical textbooks. The level of difficulty varies widely and many readers may be confused or discouraged by the mixture of all the given methods described or referred to. A more systematic arrangement of the examples demonstrated is desired. The book is not recommended as a first handbook of reliability. Anyhow, the book may be used as a reference book, giving people advanced in reliability analysis to find new ideas how to solve reliability problems.

The first three chapters treating fundamental concepts of reliability, e.g. some basic probability distributions, censoring of life data and maximum likelihood (ML) methods of point estimation are the most accessible to nonspecialists.

However, the 3 parameter Weibull and the bimodal Weibull distributions should have been more extensively treated, e.g. the estimation problems of the guarantee and the proportion parameters. The modified maximum likelihood estimators (MMLE) is not mentioned (Cohen *Technometrics* V. 17 1975) and the proportion parameter ( $p$ ) of the "weak" population of a bimodal distribution, for example, is easily estimated by a graphical method. The important concepts of hypothesis testing and confidence intervals are not sufficiently explained and exemplified. For graphical plotting, the authors do not mention the difference between the plotting positions representing the mean  $\log\{-\log[1-(i-1/2)/n]\}$  and the median  $\log\{-\log[1-(i-0.3)/(n+0.4)]\}$ . The last plotting position being the most common in industrial applications.

Furthermore, it is well-known that the ML-method gives biased point estimates, especially at small test sizes ( $n$ ) and/or strong censoring ( $r/n$ ). All ML point estimates of the examples reported are "raw" values uncorrected for bias. For example, in example 2.2, assuming a 2-parameter Weibull distribution, the observed "raw" ML-estimates of (Weibull slope) and  $\alpha$  (Characteristic life) are 1.42 and 2.27, whilst the median bias corrected values are 1.26 and

2.40, respectively, as obtained by extensive Monte Carlo simulated experiments, when  $n=13$  and  $r=10$ . Now, the point estimates alone are virtually useless in reliability engineering. It is of more value to have a measure of the precision or confidence (reliability) in the observed estimates, e.g. by the estimated confidence intervals. For example, the 90% confidence intervals of the two Weibull parameters of example 2.2 are (0.73,1.97) and (1.59,4.03), respectively.

Chapter 4 describes regression models for reliability data following the Weibull model and taking other covariate or explanatory variables into account. Accelerated life models are treated by several models, as proportional hazards or more generalized models. The common methods of the power rule model (Levenbach 1957), the Arrhenius reaction rate model (Thomas and Gorton 1963, Pershing 1964) and the generalised Eyring model (Goldberg 1964) are missing. By accelerated test the change of the observed failure modes is also of extreme importance and must seriously be taken into account. By and large, many important and valuable articles (authors) are missing in the given bibliography over the first five chapters.

Chapter 5 is dealing with proportional hazards (PH) modelling. This chapter may be difficult to understand and apply for an engineer looking for a convincing solution and analysis of his observed life data. Sometimes, more straightforward methods may give simpler, but equally realistic and reliable, solutions.

Chapter 6 deals with Bayesian statistics. The theory of the Bayesian statistics is clearly treated. However, the Bayesian statistics in reliability is not often met in practise. The difficulty of the Bayesian approach lies in the specification (by experience or historical data) of the prior distribution of the unknown (random) distribution parameters in order to determine the posterior distribution and the predictive survivor function.

The Gamma(a,b) distribution, belonging to the natural conjugate family, is exemplified, for the prior distribution of the MTTF of the exponential distribution. Bayesian decision analysis, e.g. of safe inspection intervals or optimal replacements, is demonstrated.

Chapter 7 deals shortly with multivariate models. Here, knowledge and access of MANOVA (multivariate analysis of variance) program is needed.

The two last chapters are dealing with the reliability of repairable systems and models for system reliability.

The chapters (only 76 pages) gives clear definitions of reliability concepts as ROCOF (rate of occurrence of failures), non-homogeneous Poisson process models, repairable and non-repairable systems, coherent systems and structure and reliable functions. However, the chapters are demanding an advanced level of statistical knowledge of the reader.

Summarised, the book is readable and of interest for the professional statisticians (experts) in the reliability area. Some chapters may also be readable and of general use for engineers having a good statistical knowledge.



The content of the book is not easy to digest, but the book may anyhow be useful as a general reference source for the statistical libraries.

Thyr Andersson,

Consulting Statistician, SKF Engineering & Research Centre, B.V.

### **E.K. Foreman**

#### *Survey sampling principles*

Marcel Dekker, New York, 1991, xvi + 497, ISBN 0-8247-8407-3, \$ 126.50.

E.K. Foreman's boek over steekproeftheorie heeft een lage "return on investment". De opgedane kennis en inzichten steken magertjes af tegen de prijs (\$ 126.50), maar meer tegen de geïnvesteerde tijd, nodig om de 497 bladzijden door te nemen.

Waarom zo negatief over dit boek ? Daar wil ik drie redenen voor aangeven. Ten eerste is de auteur erg breedsprakig. Grote hoeveelheden tekst worden besteed aan onderwerpen zoals : "Control of nonsample errors", dat na lezing niet meer oplevert dan het besef dat je erg je best moet doen om die fouten onder controle te houden. Ten tweede worden de mogelijkheden van het gebruik van modellen voor het schatten en toetsen van populatieparameters onderbelicht. Dat vermoeden deed ik reeds op met de openingszin van het boek: "Sampling is an example of inductive logic by which conclusions are inferred on the basis of a limited number of instances". Naar mijn idee horen modellen en toetsen meer bij deductie.

Bij de stukken waar het gaat over ratio- en regressieschatters wordt het regressiemodel te laat gebruikt en wordt de regressor te lang alleen als "benchmark" gezien. Dan blijkt een intercept ( $\alpha$  in de vergelijking  $y = \alpha + \beta x$ ) maar lastig: "The constant value  $\alpha$  of the intercept could be estimated and deducted from each of the benchmark values  $X_1$ . (blz.86), of : "When the intercept  $\alpha$  is negligible, as often happens in practice... (blz 85).

Bovendien lijkt het multipele regressie niet gewenst, want er wordt aangeraden te kiezen uit een groep regressoren (blz. 95) en later heet het "complexe regressie" (blz. 264). Verder wordt nergens het superpopulatiemodel behandeld. In een superpopulatiemodel wordt de onderhavige populatie gezien als een steekproef uit een "super"-populatie en levert een handig kader om eigenschappen van allerlei schatters af te leiden.

Ten derde is de mathematische behandeling is zelden fout maar niet exact. Er wordt in de notatie geen onderscheid gemaakt tussen variabelen voor en na de trekking. En wat te denken van een zin als: "This assumption - of a sample of equally probable elements representing an infinite population - ... (blz 2). En een duistere zin als " ...where  $E^*$  denotes an expectation of a

special kind." Bovendien worden belangrijke beweringen niet of nauwelijks mathematisch onderbouwd; de eindigheidscorrectie komt compleet uit de lucht vallen.

Samenvattend ik zou iemand dit boek niet snel aanraden. Voor mensen, die Nederlands kunnen lezen is de syllabus van Bethlehem (1985), uitgegeven door het CBS, een prima start, waar je vlot doorheen bent (222 blz). Het nieuwe boek van Muilwijk, Snijders en Moors (1992) lijkt mij heel goed; de recensie staat elders in dit nummer van KM. Verder blijft het (engelstalige) boek van Cochran (1973) een solide werk. Slechts voor mensen die Cochran te wiskundig vinden en veel tijd hebben is Foreman's boek een aanrader.

#### Referenties

- Bethlehem, J.C. (1985) *Theorie en praktijk van het steekproefonderzoek*. CBS, Voorburg  
 Cochran, W. (1973) *Sampling techniques*. John Wiley & Sons, New York  
 J. Muilwijk, T.A.B. Snijders en J.J.A. Moors (1992). *Kanssteekproeven*. Stenfert Kroese, Leiden

Thom Luijben  
 PTT

#### Alan Pankratz

##### *Forecasting with dynamic regression models*

John Wiley & Sons, New York, 1991, xiii + 386 pag., ISBN 0-471-61528-5, £ 47,50

Dit boek behandelt het maken van voorspellingen van één continue variabele, met behulp van ARIMA-modellen met één of meer verklarende variabelen. De nadruk ligt op de praktijk en niet op de wiskundige statistiek. Het boek is een *companion volume* bij een vorig boek van dezelfde auteur (*Forecasting with univariate Box-Jenkins models: concepts and cases*, Wiley, 1983).

Na een inleidend hoofdstuk volgt een *primer on ARIMA models*. Hierin worden ARIMA-modellen (met eventueel seizoenen) geïntroduceerd en wordt de relatie tussen het model en de autocorrelatie-coëfficiënten beschreven.

Dan volgt een *primer on regression models*. Hierin wordt (aan de hand van doorsnedegegevens) het lineaire regressiemodel behandeld, inclusief OLS en haar eigenschappen, het toetsen van een hypothese betreffende het model, enz. Tevens wordt hier de storingsterm als



ARMA-proces geïntroduceerd. De formulering van een ARMA-model middels een quotiënt van twee polynomen in de vertragingoperator komt hier het eerst voor.

In het volgende hoofdstuk wordt hieraan toegevoegd een rationale vertraging op de verklarende variabelen, geïntroduceerd via de Koyckse vertraging.

Dan begint, op bladzij 167, het eigenlijke werk: het "identificeren" van het model, in de zin van Box en Jenkins. Hierbij ook aandacht voor het testen op simultane beïnvloeding: niet alleen van  $X$  naar  $Y$ , maar ook omgekeerd. Dit hoofdstuk wordt gevolgd door een hoofdstuk waarin wordt besproken hoe een geschat model aan de tand moet worden gevoeld: de *diagnostic checks*.

In de twee volgende hoofdstukken (7 en 8) worden modellen met schoksgewijze veranderingen besproken, alsmede het opsporen van zulke veranderingen (*intervention analysis*). Hiermee hangt samen de opsporing en de behandeling van uitschieters. Deze hoofdstukken lijken een beetje samengeraapt; hoofdstuk 8 begint met het geven van andere namen aan de verschijnselen die in hoofdstuk 7 zijn besproken.

Hoofdstuk 9 gaat over het schatten van de eerder besproken modellen en het berekenen van een voorspelling daarmee. Uit het oogpunt van de wiskundige statistiek is het schatten wellicht de kern van het boek, maar - overeenkomstig het uitgangspunt van de auteur - wordt dit slechts behandeld voor zover nodig is om een idee te krijgen van de verschillen tussen de schattingsmethoden en min of meer oordeelkundig computerprogrammatuur te gebruiken. Meer in het algemeen streeft de auteur meer naar het aankweken van intuïtie bij de lezer, dan naar het streng afleiden van de benodigde stellingen.

In het laatste hoofdstuk komen vector-ARMA-modellen aan de orde. Hierbij worden de betrokken reeksen symmetrisch gemodelleerd: in beginsel hangt iedere variabele af van het eigen verleden en dat van alle andere variabelen.

Tussen de hoofdstukken in worden zes praktische problemen gedetailleerd uitgewerkt. Dit geeft een aanmerkelijke verrijking voor degene die zelf aan de slag moet. Voorts worden aan het slot van ieder hoofdstuk opgaven gesteld.

Dit lijkt me een tamelijk nuttig boek voor de lezers waarvoor het bedoeld is: zij die voorspellingen moeten maken van één of een klein aantal variabelen en daarbij kunnen volstaan met een niet-simultaan model (of het desbetreffende onderwijs). Dit ondanks de hiernavolgende minpunten.

Er wordt in het boek weinig aandacht gegeven aan het bekijken van de residuen om te zien of er nog verklarende variabelen bij moeten. Geen coïntegratie. Geen Almon-vertragingen. Bij de presentatie van resultaten zijn soms te veel decimalen achter de punt afgedrukt. (Desgevraagd verklaarde de auteur dat hij vanwege tijdsdruk hier niet voldoende aandacht aan heeft gegeven.) De formule voor de geschatte helling van een regressielijn op blz. 87 mist twee van de drie openingshaken.

Tenslotte: in het uitgebreide rekenvoorbeeld in hoofdstuk 9 is bij de bepaling van standaardfout en betrouwbaarheidsinterval rond de voorspelling uitgegaan van een onzekere verklarende variabele in de toekomst. Echter, de voorspelling is vermoedelijk uitgevoerd met de (in dit geval) over het voorspeltijdvak bekende waarde van de verklarende variabele. Dit levert een wonderlijke figuur 9.1 op, met een breed uitwaaiend interval en daarbinnen een vrijwel foutloze voorspelling. (De echte waarde van de afhankelijke variabele is ook in de figuur getekend.)

Arie ten Cate

Centraal Planbureau, Den Haag

### **T.P. Hutchinson**

*Ability, partial information, guessing: statistical modelling applied to multiple-choice tests*

Rumsby Scientific Publishing, Adelaide, 1991. xiv + 266 pag., ISBN 0-646-03328-X, \$ 24.00.

Al meer dan zeventig jaar wordt de mensheid geplaagd met multiple-choice-tests. Van de wieg tot het graf wordt men gemeten met meerkeuze-opgaven. Niettemin zijn de statistische modellen die een beschrijving geven van de psychologische processen van het "antwoordgeven" uiterst primitief. Tegelijkertijd zijn de theorieën uit de cognitieve psychologie te specialistisch om van nut te kunnen zijn in de praktijk. Het besproken boek heeft de pretentie, aan deze misstanden een einde te maken. Daartoe wordt er een nieuwe theorie ontvouwd: de "past-niet-theorie" (mismatch theory). En passant wordt men onderwezen in de itemresponsiemodellen, het effect van verschillende vraagvormen en het scoren (met punten waarderen) van een antwoord.

Bovenstaande alinea is een parafrase van het eerste hoofdstuk van het te bespreken boek. Aan het ambitieuze programma voegt de schrijver vernietigende kritiek toe op de itemresponsie-modellen uit de psychometrie, waarmee hij het boek tot een pamflet maakt. Is de schrijver geslaagd in zijn opzet, een gulden middenweg te vinden tussen de barre woestijn van de psychometrie en de droomwereld van de cognitieve psychologie? Laat ons, voordat we deze vraag beantwoorden, eens kijken naar de "past-niet-theorie".

De "past-niet-theorie" is van toepassing op meerkeuzevragen die  $N$  antwoordmogelijkheden hebben waarvan er een het correcte antwoord is. Volgens de theorie trekt iemand die met een opgave wordt geconfronteerd, voor elk der antwoordmogelijkheden aselect uit een verdeling. Als de trekking een drempelwaarde  $T$  overschrijdt, wordt de antwoordmogelijkheid in kwestie *niet* aangekruist. Voor het correcte antwoord wordt er getrokken uit de verdeling  $F$ , voor alle andere uit een en dezelfde verdeling  $G$ . Er doen zich nu drie gevallen voor:



- de kleinste van de N getrokken waarden is groter dan T: er wordt geen enkele antwoordmogelijkheid aangekruist (de opgave wordt overgeslagen);
- de kleinste van de N getrokken waarden is niet groter dan T en afkomstig uit G: er wordt een onjuist antwoord aangekruist;
- de kleinste van de N getrokken waarden is niet groter dan T en afkomstig uit F: het correcte antwoord wordt aangekruist.

Een belangrijke assumptie van de theorie is dat de verdeling F stochastisch kleiner is dan G. De "afstand" tussen F en G is een niet-dalende functie van een parameter  $\lambda$ . Hoe groter  $\lambda$ , des te groter is de kans op een correct antwoord.

De notie dat een antwoord op een opgave het resultaat is van een trekking uit een verdeling, gaat terug op Thurstone (1927). Het is een vruchtbaar idee gebleken, dat tot een veelheid aan psychometrische meetmodellen heeft geleid. Het is aardig, te weten dat Thurstone het trekken uit de verdelingen F en G opzettelijk aanduidde met de betekenisloze term "psychologisch proces" om te benadrukken dat het hem ging om de procedure die tot een waarneembare uitkomst leidt. De "past-niet-theorie" sluit hier naadloos bij aan. In het boek wordt geen poging ondernomen, het "psychologische proces" substantie te geven. Daarmee is een van de in het eerste hoofdstuk genoemde pretenties niet waargemaakt.

In het boek worden veel voorbeelden gegeven van verdelingen F en G die aan de veronderstellingen van de theorie voldoen. Hoe men de verdelingen moet kiezen, is een raadsel. In diverse voorbeelden worden allerlei varianten losgelaten op dezelfde gegevens, en op de een of andere wijze voor al dan niet passend verklaard.

Merk op dat er in de theorie slechts een enkele parameter  $\lambda$  optreedt. Deze kenmerkt de confrontatie van een subject met een opgave. De parameter heeft geen subscript; de I•J confrontaties van I subjecten met J opgaven worden beschouwd als I•J replicaties. Dat leidt tot een uitermate eenvoudige schattingsmethode. Men bepaalt de proporties overgeslagen, correct beantwoorde en incorrect beantwoorde opgaven in de I•J antwoorden. Deze proporties worden beschouwd als de kansen op elk der drie soorten antwoorden. De theorie beschrijft de kansen als functies van de drempelwaarde T en de parameter  $\lambda$ . Deze kan men eenvoudig uit de proporties berekenen.

Of  $\lambda$  de moeilijkheidsgraad van de opgaven dan wel het vaardigheidsniveau van de subjecten karakteriseert, is niet uit te maken. Daarmee negeert de "past-niet-theorie" de ontwikkelingen van de laatste dertig tot veertig jaar. De psychometrie heeft in die tijd modellen ontwikkeld waarin het in beginsel mogelijk is de karakteristieken van subjecten en opgaven van elkaar te scheiden; deze modellen staan bekend als *itemresponsi modellen*. Voor de schrijver is de theorie aangaande itemresponsi modellen "a grotesque monster" (bladz. 16). Aangezien de "past-niet-theorie" niets meer is dan een armzalig itemresponsi model, zoals ik daarnet heb uiteengezet, maakt de schrijver zich met deze uitspraak nogal belachelijk. Bovendien schroomt hij niet, zich in hoofdstuk 9 met het gedrocht van de itemresponsi modellen te encanaileren. In dat hoofdstuk

opert hij diverse mogelijkheden om de parameter  $\lambda$ , nu plotseling voorzien van subscripts  $i$  en  $j$ , te ontbinden in een subjectparameter  $\vartheta_i$  en een parameter  $b_j$  voor de opgave! Deze ontbindingen leiden vanzelf tot een aantal bekende itemresponsi modellen. Het zou de moeite waard zijn, zo lezen we in hoofdstuk 9, om voor deze modellen schattingsmethoden te ontwikkelen. Wie is hier nu "grotesque"?

Tussen hoofdstuk 3, waarin de "past-niet-theorie" wordt ontvouwd, en hoofdstuk 9 waar itemresponsi modellen tot bijzondere gevallen van deze theorie worden verklaard, bevinden zich vijf hoofdstukken waarin de schrijver diverse vraagvormen en soorten instructies beschrijft. Als men de toepassingen van de "past-niet-theorie" negeert, heeft men in deze hoofdstukken een aardige inventaris van vraagvormen en soorten instructies.

Het boek van Hutchinson is een raar boek. Niet geplaagd door bescheidenheid en kennis van zaken, gaat de schrijver tekeer tegen wat de psychometrie de afgelopen veertig jaar heeft opgeleverd. De pretenties die hij heeft, worden echter op geen enkele manier waargemaakt. Een minder hoog ambitieniveau en meer ingetogenheid hadden tot een leuk boek kunnen leiden.

### Referentie

L.L. Thurstone (1927). A Law of Comparative Judgement. *Psychol. Rev.* 34, 237.

Niels H. Veldhuijzen  
Cito, afdeling OPD

### T.P. Hutchinson

#### *Controversies in Item Response Theory*

Rumsby Scientific Publishing, Adelaide, 1991, ii + 60 pag., ISBN 0-646-03329-8, \$ 7.00.

Dit boekje in A6-formaat is een parafrase van het tweede, en een deel van het achtste, hoofdstuk van het hierboven beschreven boek. De laatste alinea van bovenstaande recensie is evenzeer van toepassing op dit kleinood.

Niels H. Veldhuijzen  
Cito, afdeling OPD



Lyman Ott and William Mendenhall

*Understanding statistics*

PWS-Kent, Boston, 1990, xvii + 727 pag., ISBN 0-534-98169-0, £ 15.95.

Dit boek is bedoeld als inleidende cursus in de statistiek, voor studenten in de economische, sociale en natuurwetenschappen. Het boek is vlot geschreven met veel voorbeelden en wat dat betreft is het typisch Amerikaans. Het accent wordt gelegd op statistiek als "making sense of data". De hoeveelheid uitleg is in het algemeen ruim, zonder op de knieën te gaan zitten. Het boek bevat opgaven uit diverse hoek, waarvan sommige met uitwerkingen. Daarnaast zijn ook software voorbeelden uit SAS en Minitab. Ook voor zelfstudie lijkt het boek goed te gebruiken.

In vergelijking met de gebruikelijke opzet van een leerboek in de statistiek is er in dit boek wat meer aandacht voor het verzamelen van gegevens. Na enkele inleidende hoofdstukken hierover volgen twee hoofdstukken over beschrijvende statistiek, met veel aandacht voor grafische presentatie. Daarna volgt een duidelijke bespreking van kansverdelingen. De auteurs behandelen de normale en binomiale verdeling en de bijbehorende twee-steekproefproblemen. De Poisson-verdeling ontbreekt helaas.

Het boek gaat dan verder met regressieanalyse en een inleiding in proefopzetten. De behandeling van correlatiecoëfficiënt is hier wat zwakjes.

Verder komt variantieanalyse aan bod. Een korte, maar heldere inleiding. De auteurs beperken zich tot een-weg ANOVA maar gaan wel uit van het geval van ongelijke steekproefgrootten. De argeloze student wordt gelukkig niet bestookt met "multiple comparison"-toetsen.

Bij de behandeling van het twee-steekproefprobleem en variantieanalyse wordt ook ingegaan op de aannamen die gemaakt moeten worden en hoe deze getoetst kunnen worden. Voor het geval dat niet aan de voorwaarden is voldaan komt het boek helaas niet verder dan het advies "raadpleeg een statisticus". Dit soort adviezen houdt u en mij van de straat, maar een iets uitgebreidere behandeling van bijvoorbeeld transformaties was toch wel op zijn plaats geweest.

Aan het eind van het boek wordt nog kort ingegaan op niet-parametrische toetsen. Een beetje een verplicht nummertje. Het hele boek staat in het teken van de parametrische statistiek en zo'n hoofdstuk wordt dan meestal toch overgeslagen. Nuttig is wel een hoofdstuk over het communiceren van het resultaat van statistische analyses. Dit blijft echter enigszins hangen in een opsomming van wat er allemaal fout kan gaan in plaats van de student iets te willen leren over dit onderwerp.

Samenvattend: de auteurs zijn er in geslaagd een aantal statistische onderwerpen op heldere wijze te presenteren. Qua inhoud is het boek mijns inziens echter te weinig uitgebalanceerd om een inleidende cursus helemaal aan op te hangen - er zijn immers voldoende alternatieven. Als bron van aanvullend materiaal niettemin zeker aanbevolen.

>>>>>>>>>> **Binnengekomen boeken** <<<<<<<<<<<

Om de nieuws waarde van de boekbesprekingsrubriek te verhogen wordt een lijst gepresenteerd van boeken die in de afgelopen periode bij de redactie zijn binnengekomen. Omdat er soms nogal wat tijd kan zitten tussen de uitgave van het boek en de publicatie van de recensie, hoopt de redactie hiermee de dienstverlening aan de lezers te vergroten.

Niet alle boeken in deze rubriek zijn al ondergebracht bij een recensent. Mocht er belangstelling bestaan voor het bespreken van een boek, dan kan contact opgenomen worden met de boekbesprekingsredacteur. Met nadruk wordt gesteld dat niet zeker is of het boek nog 'in de aanbieding' is.

**LePage, R. & L. Billard**

*Exploring the limits of bootstrap*

John Wiley & Sons, New York, 1992, xi + 426 pag., ISBN 0-471-53631-8, £ 47.50

**Montgomery, D.C. & E.A. Peck**

*Introduction to linear regression analysis, 2nd edition*

John Wiley & Sons, New York, 1992, xiv + 527 pag., ISBN 0-471-53387-4, £ 39.95

**S.R. Searle, G. Caselle & C.E. McCullach**

*Variance components*

John Wiley & Sons, New York, 1992, xxiii + 501 pag., ISBN 0-471-62162-5, £ 56.00

**Box, G.E.P. & G.C. Tiao**

*Bayesian inference in statistical analysis*

John Wiley & Sons, New York, 1992, xviii + 588 pag., ISBN 0-471-57428-7, £ 30.95

**Lessler, J.T. & W.D. Kalsbeek**

*Nonsampling errors in surveys*

John Wiley & Sons, New York, 1992, xiii + 412 pag., ISBN 0-471-86908-2, £ 56.00

**Harn, K. van & P.J. Holeyijn**

*Markov-ketens in diskrete tijd*

Epsilon uitgaven, Utrecht, 1991, v + 244 pag., ISBN 90-5041-026-X, f 34.50

**Choi, B.S.**

*ARMA model identification*

Springer-Verlag, Berlijn, 1992, xi + 200 pag., ISBN 3-540-97795-3, DM 78.00



**Chase, W. & F. Bown**

*General Statistics, second edition*

John Wiley & Sons, New York, 1992, xiv + 720 pag., ISBN 0-471-55915-6, £ 18.95

> > > > > > > **Proefschriften, CWI uitgaven** < < < < < < < <

**F. Thuijsman**

*Optimality and equilibria in stochastic games*

Stichting Mathematisch Centrum, Amsterdam, 1992, iii + 107 pag., ISBN 90-6196-406-7, f 40.00.

**Koopman, S.J.**

*Diagnostic checking and Intra-daily effects in Time series models*

Thesis Publishers, Amsterdam, 1992, xiii + 236 pag., ISBN 90-5170-142-X, prijs onbekend

Dit is het proefschrift van de auteur op basis waarvan hij op 24 april 1992 is gepromoveerd aan de London School of Economics and Political Science te Londen. Het proefschrift is verschenen in de Tinbergen Institute Research Series en kan worden besteld via Thesis Publishers, Postbus 14791, 1001 LG Amsterdam.

## Oproep voor boekbespekers

Als redacteur van de boekbesprekingsrubriek van de VVS, welke gepubliceerd wordt in Kwantitatieve Methoden, ben ik steeds op zoek naar mensen die bereid zijn een pas verschenen boek te willen lezen met het doel daarover een oordeel te geven. Dat oordeel wordt op prijs gesteld door de leden van de VVS. Het verschaft inzicht in de kwaliteiten van het boek. Voor de schrijver(s) kan een oordeel van het publiek waarvoor het boek geschreven is een mogelijkheid bieden tot verdere verbetering. Ook de uitgever heeft er uiteraard baat bij. Als tegenprestatie voor het beoordelen mogen de recensenten de besproken boeken behouden.

Behalve de algemene boeken over statistiek, waar veel mensen zinnige dingen over kunnen zeggen, betreft het vaak ook boeken met een specialistisch onderwerp. Als ik, soms op de gok, iemand een dergelijk boek toestuur, dan hoor ik een enkele keer als reactie dat 'het boek precies aansluit bij het onderzoek' dat de recensent verricht. 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, wil ik graag in contact komen met veel collega statistici, kansrekenaars en OR-beoefenaren, al zijn de boeken op dit laatste gebied schaars. 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 boeken binnenkomen heb ik een grotere keuze van mogelijke recensenten en u wellicht een gloednieuw boek op uw onderzoeksgebied.

Ja, ik ben graag bereid een boek op mijn interessegebied te bespreken voor de leden van de VVS

naam

-----

adres

-----

-----

-----

-----

interessegebied (graag enigszins nauwkeurig  
omschrijven)

-----

-----

-----