

Boekbesprekingen

Redactie

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

telefoon 013 - 640319

Besproken boeken in Kwantitatieve Methoden nr. 42

Joe Whittaker	179
<i>Graphical Models in Applied Multivariate Analysis</i>	
A.J. Lee	181
<i>U-Statistics: Theory and Practice</i>	
Michel Emery	182
<i>Stochastic Calculus in Manifolds</i>	
A. Buja en Paul A. Tukey, editors	183
<i>Computing and Graphics in Statistics</i>	
J.S. Maritz and T. Lwin	187
<i>Empirical Bayes Methods, 2nd edition</i>	
A. Brandt, P. Franken and B. Lisek	189
<i>Stationary Stochastic Models</i>	
G.E.P. Box and G.C. Tiao	191
<i>Bayesian Inference in Statistical Analysis</i>	
D.M. Eddy, V. Hasselblad, and R. Shachter.	193
<i>Meta-Analysis by the Confidence Profile Method: The Statistical Synthesis of Evidence</i>	
D.C. Montgomery en E.A. Peck	195
<i>Introduction to linear regression analysis</i>	
D.C. Hoaglin, F. Mosteller & J.W. Tukey	196
<i>Fundamentals of exploratory analysis of variance</i>	
H. Leon Harter	199
<i>The chronological annotated bibliography of order statistics, vol VI: 1966-1967</i>	
H. Leon Harter	199
<i>The chronological annotated bibliography of order statistics, vol VII: 1968-1969</i>	

Joe Whittaker

Graphical Models in Applied Multivariate Analysis

John Wiley & Sons, New York, 1990, xiv + 448 pag., ISBN 0-471-91750-8, £ 39.95

Het uitgangspunt van "graphical modelling" is dat modellen voor conditionele onafhankelijkheid in multivariate verdelingen als een graaf kunnen worden weergegeven. Men mikt dus niet, zoals in de meer conventionele aanpak van multivariate analyse, op een beschrijving van onderlinge afhankelijkheden, maar op een beschrijving van onafhankelijkheden, Whittaker's boek geeft als eerste een uitgebreid overzicht van dit zich snel ontwikkelende gebied, dat tracht unificerende concepten aan te reiken voor de gezamenlijke behandeling van discrete en continue variabelen, daarmee een brug slaand tussen apart ontwikkelde methoden voor de analyse van contingentie tabellen enerzijds en correlatie matrices anderzijds.

Een centrale plaats in het boek wordt ingenomen door Markov eigenschappen van een onafhankelijksgraaf, in navolging van artikelen zoals Darroch, Lauritzen & Speed (1980), Kiiviri & Speed (1982), en Lauritzen & Wermuth (1989), en het geeft naast de uitleg van de theorie een groot aantal van deze beschouwingswijze. De hoofdstuk indeling is als volgt:

- 1 Introduction
- 2 Independence Interaction
- 3 Independence Graphs
- 4 Information Divergence
- 5 The Inverse Variance
- 6 Graphical Gaussian Models
- 7 Graphical Log-linear Models
- 8 Model Selection
- 9 Methods for Sparse Tables
- 10 Regression and Graphical Chain models
- 11 Models for Mixed Variables
- 12 Decompositions and Decomposability

Systemen van conditionele onafhankelijkheid zijn niet het enige raamwerk waarin een geïntegreerde behandeling van numerieke en categorische variabelen mogelijk is. Een groot deel van de huidige toegepaste meerdimensionale data analyse maakt gebruik van een gewone vectorruimte, waarin variabelen gerepresenteerd zijn als vectoren of groepen van vectoren. Deze wijze van representeren heeft als voordeel dat het eenvoudiger is om een groep waargenomen variabelen (of de categorieën van een variabele) als een steekproef uit een domein van mogelijke variabelen (of categorieën te zien. Het domein (de "populatie" van variabelen/categorieën) kan een deelruimte zijn, die wordt opgespannen door een beperkt aantal vectoren (latente variabelen), waarbinnen elke steekproef van actuele

variabelen zich zou bevinden. In graphical modelling wordt daarentegen uitgegaan van een vaste groep van variabelen, en er wordt weinig speciale aandacht gegeven aan het idee van de latente variabele. Hoewel Spearman's (1904) klassiek één factor-model toch meestal geformuleerd wordt als onafhankelijkheid van ieder paar intelligentie test, conditioneel op de latente variabele algemene intelligentie. Dit basismodel is een ster-vormige graaf, met de latente variabele in het centrum. Wellicht valt er binnen het kader van de grafische modellen theorie verder weinig over te zeggen.

Ook andere, meer recente modellen in de biometrie en psychometrie voor het onderscheiden van subpopulaties gaan uit van wat dan genoemd wordt lokale onafhankelijkheid, d.w.z. opnieuw onafhankelijkheid van waargenomen responsen, gegeven de lokatie van de (bekende of onbekende) subpopulatie op één of meer latente variabele(n). Een overzicht hiervan wordt gegeven in Langeheine & Rost (1988). Heeft het zin ze te beschouwen als grafische modellen? We weten het niet, want deze literatuur is helaas geheel aan Whittaker voorbij gegaan. En dat terwijl men toch ook de "modelaanpak" in de multivariate analyse nastreeft en ruim de nadruk legt op onafhankelijkheidsaannamen.

Deze bemerkingen weerspiegelen voornamelijk eigen preoccupaties en nemen geenszins weg dat ik het boek van harte kan aanraden als naslagwerk of als tekst voor een gevorderde cursus in de multivariate analyse. Het brengt deze toch stevige kost op zeer heldere, soms zelfs lichtvoetige wijze, met veel zijdelingse opmerkingen die het eenvoudiger maken de portee van de geciteerde artikelen te begrijpen. Het boek bevat ook een behoorlijk aantal oefeningen, voornamelijk van theoretische aard, met schetsen van de goede antwoorden. Wel is het boek voor statistici; het woord "toegepast" in de titel betekent niet dat een empirische wetenschapper er uit kan leren hoe men de multivariate analyse kan toepassen, maar is in de wiskundige betekenis bedoeld.

- Darroch, J.N., Lauritzen, S.L., & Speed, T.P. (1980). Markov fields and log linear interaction models for contingency tables. *Ann. Statist.*, 8, 522-539.
- Kiiveri, H., & Speed, T.P. (1982). Structural analysis of multivariate data: a review. In; Leinhardt, S. (Ed.), *Sociological Methodology*. San Francisco: Jossey Bass.
- Langeheine, R., & Rost, J. (1988). *Latent trait and latent Class Models*. New York en Londen: Pleunum Press.
- Lauritzen, S.L. & Wermuth, N. (1989). graphical models for associations between variables, some of which are qualitative and some quantitative. *Ann. Statist.*, 17, 31-54
- Spearman, C. (1904). General intelligence, objectively determined and measured. *Amer. J. of Psych.*, 15, 201-293.

W.J. Heiser

vakgroep datatheorie en wetenschapsstudies RUL

A.J. Lee

U-Statistics: Theory and Practice

Marcel Dekker, New York, 1990, xii + 302 pag., ISBN 0-0247-8253-4, \$ 95.70

De klasse van de U-Statistics is ongeveer 40 jaar geleden geïntroduceerd door Paul Halmos en Wassily Hoeffding. In de daaropvolgende tijd zijn er veel artikelen en delen van boeken aan deze klasse van statistics gewijd. De auteur stelt zich ten doel de theorie van de U-Statistics op een overzichtelijke manier in een boek te beschrijven en daarnaast nader in te gaan op een aantal toepassingen. Om een indruk te krijgen van de van de besproken onderwerpen geven we een kort overzicht van de hoofdstukken.

In Hoofdstuk 1, "Basics", wordt het belang van U-Statistics geïllustreerd door een bespreking van hun optimaliteitseigenschappen als schatters. Bovendien worden er uitdrukkingen afgeleid voor de momenten van U-Statistics. Het belangrijkste onderdeel is misschien wel de afleiding van de zogenaamde Hoeffding decompositie van een U-Statistic in termen van een som van ongecorreleerde U-Statistics van lagere graad. Deze decompositie is het basisresultaat dat zeer veel in het vervolg aangeroept wordt.

In Hoofdstuk 2, "Variations", worden afwijkingen van de in Hoofdstuk 1 beschouwde i.i.d. standaard situatie besproken. We noemen: gegeneraliseerde U-Statistics, afwijkingen van het i.i.d. geval, U-Statistics gebaseerd op trekkingen zonder teruglegging, gewogen U-Statistics en gegeneraliseerde L-Statistics.

In Hoofdstuk 3, "Asymptotics", worden achtereenvolgens convergentie in verdeling, ook voor gedegeneerde U-Statistics, convergentiesnelheden, de sterke wet voor U-Statistics, de wet van de geïteerde logaritme en invariantie principes besproken. Daarna worden de asymptotiek voor de statistics uit het voorafgaande hoofdstuk en de asymptotiek voor het geval van kernen met geschatte parameters behandeld. Aardig is dat er van sterke wet drie verschillende bewijzen worden gegeven. Het eerste bewijs is het originele bewijs van Hoeffding uit een ongepubliceerd technisch rapport uit 1961. De andere bewijzen gebruiken martingaal technieken en de Hewitt-Savage nul-een wet.

In Hoofdstuk 4, "Related statistics", worden algemene symmetrische statistics besproken, alsmede V-Statistics en Incomplete U-Statistics.

In Hoofdstuk 5, "Estimating standard errors", worden de Jackknife en Bootstrap gebruikt om de variantie van standaard- en Incomplete U-Statistics te schatten.

In Hoofdstuk 6, "Applications", worden U-Statistics ingevoerd voor een aantal specifieke schattings- en toetsingsproblemen. We noemen zonder volledig te zijn: circulaire en sferische correlatie, toetsen van symmetrie, toetsen van normaliteit, toetsen van onafhankelijkheid, meersteekproevenproblemen en een toets voor het "new better than used" probleem.

De auteur slaagt er inderdaad in een overzicht te geven van de huidige stand van zaken. De bewijzen die gegeven worden komen doorgaans uit recente artikelen, zij het dat hier en daar ten gunste van de presentatie eenvoudigere gevallen bekeken worden. Er worden

redelijk veel voorbeelden uitgewerkt om de theorie te illustreren. Waar wat lastigere theorie wordt aangeroepen, zoals bijvoorbeeld convergentie naar het Wiener proces, of martingalen, geeft de auteur een korte samenvatting van deze theorie, alvorens die toe te passen. Het is een nuttig boek voor mensen die geen of slechts gedeeltelijke kennis van de theorie van de U-Statistics.

A.J. van ES

Faculteit Wiskunde en Informatica, Universiteit van Amsterdam

Michel Emery

Stochastic Calculus in Manifolds

Springer-Verlag, Berlin, 1989, x + 151 pag., ISBN 3-540-51664-6, DM 48.00

Stochastische differentiaal meetkunde is een onderwerp dat steeds meer in de belangstelling staat. Ook in de mathematische statistiek is in het recente verleden profijt getrokken van een meetkundige benadering (zie bijvoorbeeld het werk van O. Barndorff-Nielsen). Over deze statistische toepassingen wordt in dit boek echter niet gerept. Daarin gaat het om de stochastische analyse zelf. Preciezer gezegd: het boek geeft een inleiding tot het gebruik van de zogenaamde Schwartz-Meyer tweede-orde meetkunde in de stochastische analyse.

De opzet is als volgt: In de eerste twee hoofdstukken worden definities en feiten opgesomd uit de theorie van stochastische processen en de gewone differentiaal meetkunde. Vervolgens diept Emery in de hoofdstukken 3 tot en met 5 twee belangrijke voorbeelden uit: in hoofdstuk 3 worden semi-martingalen in variëteiten gedefinieerd en de bijbehorende kwadratische variatie. Hierbij wordt meteen duidelijk dat het overbrengen van de definities uit de reële stochastische analyse op die in variëteiten additionele meetkundige structuren vereist. In het geval van de kwadratische variatie bijvoorbeeld moet als vervanger van de vermenigvuldiging van reële getallen een willekeurige bilineaire vorm worden genomen; bij de definitie van martingalen moet men op de variëteit een connectie introduceren. Dit laatste wordt in hoofdstuk 4 behandeld. Hoofdstuk 5 is een inleiding tot stochastische differentiaal meetkunde op Riemannse variëteiten (variëteiten met een "inproduct") en Brownse beweging daarop.

Pas in de laatste drie hoofdstukken wordt de algemene theorie behandeld: stochastische differentiaal vergelijkingen op variëteiten in hoofdstuk 6. Stratonovich en Ito integralen (van 1-vormen t.o.v. stochastische processen) in hoofdstuk 7 en parallel transport in hoofdstuk 8.

Tenslotte is er nog een Appendix, waarin P.A. Meyer een zeer goed leesbare uiteenzetting geeft van de theorie van semi-martingalen en stochastische analyse voor "leken".

Hoewel Emery geen voorkennis van de differentiaal meetkunde voorondersteld, lijkt mij dit

boek makkelijker leesbaar voor iemand die al een beetje vertrouwd is met de terminologie van dat vak.

Alle gebruikte definities staan er wel in, maar aangezien het boek een index is ontzegd kost het terugvinden van welk begrip dan ook onevenredig veel energie en irritatie...

Voor het overige is het boek zeer prettig geschreven: er wordt absoluut niets formeler gepresenteerd dan nodig: vaak wordt eerst een intuïtieve inleiding tot een nieuw begrip of resultaat gegeven: overal wordt duidelijk aangegeven welke aannames essentieel zijn en waarom. Verder zijn er veel mooie, uitgewerkte voorbeelden in de tekst opgenomen, maar ook opgaven die hetzij een begrip verhelderen hetzij een nieuw resultaat geven. Tenslotte is aan het eind van elk hoofdstuk een (zeer) beknopt historisch en bibliografisch overzicht te vinden.

Een mooi boek dus, en -gezien de lage prijs van boeken in de Universitext serie- ook geschikt voor een werkgroep voor studenten met (toch) de nodige vooropleiding.

Annoesjka Cabo
CWI

A. Buja en Paul A. Tukey, editors

Computing and Graphics in Statistics

Springer Verlag, Berlijn, 1991, ISBN 3-540-97633-7, xiv + 279 pag., DM 78.00

Dit boek bevat een verzameling artikelen werden die gepresenteerd in de laatste twee weken van het IMA (Institute for Mathematics and Applications van University of Minnesota) zomerprogramma van 1989. De artikelen zijn zeer divers van inhoud. Ze gaan over een breed scala van onderwerpen op het gebied van data-analyse en computertechnieken. Sommige artikelen beschrijven heel concreet bestaande computerpakketten of werkomgevingen, terwijl andere artikelen wat zweviger van aard zijn, en eigenlijk alleen nog maar wat vage ideeën bevatten.

Door de grote verscheidenheid aan onderwerpen is het moeilijk om een algemeen oordeel te geven over het boek. Voor mensen die werkzaam zijn op het terrein van de statistische informatica, zullen er vast wel enkele aardige artikelen bij zitten. Om wat orde aan te brengen heb ik getracht de artikelen enigszins op onderwerp te rangschikken, en vervolgens van elk artikel een korte karakterisering te geven.

Werkomgevingen op de computer

Een groep van zeven artikelen gaat in op statistische pakketten, of wat modieuser 'statistische werkomgevingen'. De eerste vijf artikelen beschrijven concrete systemen; de laatste twee zijn wat abstracter.

J. Cabrera: *An interface between S and Mathematica*

Dit programma beschrijft hoe men van uit het pakket S gebruik kan maken van de krachtige mogelijkheden van symbolische formule-manipulatie van het pakket Mathematica. Aldus kunnen de sterke kanten van beide pakketten worden gecombineerd.

J.A. Nelder: *GLIMPSE, a knowledge-based front end for GLIM*

GLIM is een vrij algemeen gebruikt pakket voor de analyse van gegeneraliseerde lineaire modellen. GLIMPSE kan worden opgevat als een gebruiksvriendelijke schil rondom GLIM. GLIMPSE bestaat uit drie onderdelen: een hoog-niveau stuurtaal, uitgebreide hulpfaciliteiten en een vertaler die de opdrachten vertaald in normale GLIM-opdrachten.

J.A.Nelder & R.W. Payne: *GENSTAT as a computing environment*

Statistische pakketten maken het gebruik van standaard-technieken eenvoudig, maar moedigen wat meer avontuurlijk gebruik meestal niet aan. Dit artikel beschrijft hoe men bij GENSTAT heeft gepoogd de gebruikers aan te moedigen tot meer flexibiliteit in het uitvoeren van hun analyses.

W. Dumouchel & F. O'Brien: *Integrating a robust option into a multiple regression computing environment*

Mulreg is een interactieve omgeving voor het uitvoeren van regressie-analyse. Het programma maakt deel uit van het RS/Explore pakket. Mulreg bevat veel grafische technieken die het modelleringsproces stap voor stap begeleiden. Het artikel beschrijft met name hoe een robuuste schattingstechniek (Tukey's bisquare schatter) in dit systeem is geïntegreerd.

F.W. Young & J.B. Smith: *Towards a structured data analysis environment: a cognition-based design*

De auteurs beschrijven een visueel geïoriënteerde werkomgeving voor het op gestructureerde wijze beschrijven van statistische analyses (MIDAS). Dit systeem voert niet zelf de analyse uit maar vertaalt de opdrachten in de taal van een eraan gekoppeld statistisch pakket. MIDAS gaat er vanuit dat statistici diverse cognitieve wijzen van werken hebben. MIDAS ondersteunt dat door het bieden van verschillende manieren van werken. Zo kan de statisticus zijn onderzoek grafisch beschrijven met 'data flow diagrams', maar ook kan hij textueel te werk gaan.

C.B. Hurley & R.W. Oldford: *A software model for statistical graphics*

Dit artikel beschijft een software-model voor grafische analyse. De auteur volgt daarbij de principes van de object-geïoriënteerde aanpak die tegenwoordig ook is terug te vinden bij programmeertalen als Pascal (Turbo Pascal van versie 5.5 en Turbo/Borland C++). Grafieken zijn 'views' op de gegevens, en elke view is opgebouwd uit een aantal bouwstenen die op hun beurt ook weer views kunnen zijn.

J. Pedersen: *Situations, Summaries and Model Objects*

Ook dit artikel beschrijft data-analyse vanuit een object-geïoriënteerde benadering. In dit geval betreft het IDL, een in LISP geschreven werkomgeving. Er zijn objecten die de data en meta-data beschrijven, verdelings- en transformatie-objecten beschrijven de gegevens op een wat hoger abstractie-niveau, situatie-objecten specificeren een analyse-traject en de analyses zelf zijn model- of samenvattingsobjecten.

Data-analyse

Een vijftal artikelen begeeft zich op het terrein van data-analyse:

D.B. Carr: *Looking at large data sets using binned data plots*

Wanneer een data set heel erg groot is, kan het moeilijk worden om een eventuele aanwezige structuur te ontwarren, bijvoorbeeld omdat een puntenwolk teveel punten bevat. Door de vele bomen is het bos onzichtbaar geworden. Carr beschrijft een techniek waarbij de waarnemingsruimte in zeshoeken (hexagonen) wordt verdeeld, en alle punten binnen zo'n hexagoon als één zeshoekig punt worden weergegeven, waarbij de oppervakte van dat punt het aantal weergegeven punten weerspiegelt.

J.J. Miller & E.J. Wegman: *Construction of line densities for parallel coordinate plots*

De auteurs beschrijven een techniek waarbij hoger dimensionale gegevens in het platte vlak worden weergegeven door ze af te beelden op een reeks parallelle coördinaten die elk één van de dimensies voorstellen. Om problemen met de interpretatie ten gevolge van het 'overplotten' van punten te voorkomen, werken ze met dichtheden in plaats van de punten zelf.

D.W. Scott: *On Estimation and Visualization of Higher Dimensional Surfaces*

Scott gaat in op het probleem van het visualiseren en begrijpen van structuren in hoger dimensionale data. Wat statistici meestal doen, is proberen de data in een lager dimensionale ruimte weer te geven (bijvoorbeeld met principale componenten-analyse) en daarna pas gaan ze over tot interpretatie. De auteur bespreekt enkele analyse-technieken die een zinvolle bijdragen kunnen leveren aan het oplossen van deze 'curse of dimensionality'.

W. Stuetzle: *Odds plots: A Graphical Aid for Finding Associations between Views of a Data Set*

We kunnen hoger dimensionale data bestuderen daar voor allerlei combinaties van twee variabelen naar de bijbehorende puntenwolk te kijken. Door gebruik te maken van 'brushing' kunnen we punten in een puntenwolk kleuren en dan nagaan hoe deze punten zich in andere puntenwolken gedragen. Naast een bespreking van de voor- en nadelen van brushing, presenteert de auteur ook nog de 'odds plot' als een additioneel hulpmiddel bij de interpretatie.

F.W. Young & P. Rheingans: *High-dimensional depth-cing for guided tours of multivariate data*

De auteurs beschrijven de 'guided tour' als techniek voor de verkenning van een hoog-dimensionale puntenwolk. Het computerscherm bevat de projectie van die puntenwolk terwijl het scherm als het ware door de puntenwolk reist. De 'guided tour' lijkt op de 'grand tour'. Bij de laatstgenoemde techniek kan de analist geen invloed uitoefenen op de tocht, maar bij de eerstgenoemde wel. De auteurs stellen ook het gebruik van lichtintensiteit voor om op het computerscherm het begrip afstand wat beter in beeld te brengen.

Diversen

Onder dit kopje rangschik ik een viertal overblijvende artikelen die zich niet makkelijk elders laten onderbrengen:

T. Hesterberg: *Importance sampling for Baysean estimation*

Dit artikel beschrijft het gebruik van 'importance sampling' in Bayesiaanse analyse. De klassieke importance sampling is een Monte Carlo-achtige techniek die wordt gebruikt voor het schatten van de verwachting van de functie van een stochastische variabele. Dit kan ook bij een Bayesiaanse aanpak worden gebruikt maar niet zonder enige aanpassingen. In het artikel worden die beschreven.

J.A. McDonald & J. Pedersen: *Geometric Abstractions for constrained optimization of layouts*

Dit artikel beschrijft technieken voor het maken van een layout op een beeldscherm waarmee de gebruiker op eenvoudige wijze ingewikkelde structuren kan definiëren. Dit kan variëren van het efficiënt zoeken in een database tot het weergeven van grafen. Er wordt daarbij gebruik gemaakt van technieken die afkomstig zijn uit de multidimensional scaling.

E.J. Wegman: *A Stochastic Approach to Load Balancing in Coarse Grain Parallel Computers*

De auteur presenteert modellen en enkele technieken voor het verdelen van de werklust indien een probleem wordt opgelost door gebruikmaking van parallele computers.

L. Wilkinson: *Algorithms for choosing the Domain and Range when plotting a function.*

Het construeren van een grafiek van een functie moet zodanig gebeuren dat deze op

zinvolle wijze kan worden geïnterpreteerd. De auteur definieert wat een 'informatieve grafiek' is en geeft enkele algoritmen die bij het maken ervan kunnen worden gebruikt.

Jelke Bethlehem
CBS

J.S. Maritz and T. Lwin

Empirical Bayes Methods, 2nd edition

Chapman and Hall, Londen, 1989, xi + 284 pag., ISBN 0-412-27760-3, £ 27.50

Bayesiaanse methoden in de statistiek gaan er vanuit dat de verdeling van een stochast X wordt gegeven door $F(x|\lambda)$ waarbij de vector van parameters λ zelf als stochast wordt opgevat. Volgt λ de verdeling G dan geldt

$$F(X) = \int F(X|\lambda) \quad (1)$$

In de "gewone" Bayes benadering wordt de verdeling G als subjectief opgevat.

Empirische Bayes methoden veronderstellen dat G een frequentie interpretatie toelaat, en dat data over onafhankelijke realisaties $\{X_{i,1} \dots X_{i,m(i)}; \lambda_i = 1, \dots, n, \text{ beschikbaar zijn, waarbij } X_{i,j}; j = 1, \dots, m(i) \text{ onafhankelijke trekkingen zijn uit de verdeling met parameter } \lambda_i. \text{ De waarden } \lambda_1, \dots, \lambda_n \text{ zijn niet waarneembaar.}$

In een typisch probleem wil men λ schatten aan de hand van een "current" waarde x van X , en wel zo, dat de verwachting van een verliesfunctie minimaal wordt. Zij $\delta(x)$ de schatter als functie van x , dan is een populaire keus het kwadratische verlies, waarvan de verwachting is:

$$W(\delta) = \iint (\delta(x) - \lambda)^2 dF(x|\lambda) dG(\lambda). \quad (2)$$

W wordt geminimaliseerd door $\delta(x) = E(\lambda | x)$.

De essentie van de empirische Bayes methoden bestaat hierin, dat vorige realisaties gebruikt worden in combinatie met (1) om de verdeling G te schatten. belangstelling in empirical Bayes methoden is vooral te danken aan de observatie van Robbins (1965)¹ dat de onvoorwaardelijke dichtheid $f_G(x)$ (voor continu variabelen) of massafunctie $p_G(x)$ (voor discrete variabelen) soms voldoende is om $E(\lambda | x)$ te bepalen. Robbins bestudeerde de veiligheid van verkeerskruispunten. Het aantal ongelukken per jaar op kruispunt i werd gemodelleerd als een Poisson proces met intensiteit λ_i . De intensiteiten zijn verdeeld over alle kruispunten in de stad volgens G . Zij x het aantal ongelukken op kruispunt i in een gegeven jaar dan geldt:

$$E(\lambda | X) = \frac{(1/x! \int \lambda^{x+1} e^{-\lambda} dG(\lambda))}{(1/x! \int \lambda^x e^{-\lambda} dG(\lambda))} = (x+1)p_G(X+1)/p_G(x) \quad (3)$$

p_G wordt geschat door de ongelukken te turven op andere kruispunten, dwz, kruispunten die niet dezelfde intensiteit behoeven te hebben als kruispunt 1. Als we alleen naar kruispunt i kijken, en als alleen deze ene waarneming beschikbaar is, dan zouden we schatten volgens de maximum likelihood methode, $\lambda = x$. De schatter in (3) heeft een vertragende factor die de fluctuaties bij de maximum likelihood schatter uitdempt.

Deze nieuwe editie van *Empirical Bayes Methods* van J.S. Maritz en T. Lwin beschrijft van empirische Bayes (EB) methoden voor puntschatting, intervalschatting en hypothese toetsing. Het eerste hoofdstuk geeft een leesbaar overzicht van EB methoden waar de geïnteresseerde buitenstaander veel aan zal hebben. De uitgebreide discussie van identificatie van mengsels van kansverdelingen en de vergelijking met de compound beslissingstheorie zijn de moeite waard. Een slothoofdstuk bespreekt toepassingen. Het boek bevat een uitgebreide literatuurlijst en indices voor auteurs en onderwerpen.

Over de leesbaarheid van dit boek in zijn geheel kan ik niet positief zijn. De binomiale, de Poisson en de normale modellen keren steeds weer in wisselende gedaante terug met verwijzing naar vorige besprekingen. De lezer heeft zo al gauw alle vingers van zijn linkerhand tussen de bladzijden. Het zou gemakkelijker zijn om van deze modellen drie aparte hoofdstukken te maken en daarin alle variaties te behandelen. De lezer zou dan meteen zien dat het eigenlijk om deze drie modellen gaat. Het aantal zetfouten is zo hoog dat men van slordigheid moet spreken. Veel standaard berekeningen zijn stilzwijgend als bekend verondersteld, waardoor studenten en niet-specialisten zich in de steek gelaten zullen voelen. Een definitie en korte bespreking van 'natural conjugacy' had niet mogen ontbreken.

Veel aandacht gaat uit naar de benadering van $\delta_n(x)$ van $\delta(x)$ verkregen door het vervangen van door een empirische verdelingsfunctie. Veel eenvoudiger methoden ontstaan wanneer G uit de familie van natural conjugates komt. De parameters van g kunnen dan worden geschat met de methode van momenten. Gezien de zeer geringe argumentatie om de besluiten dat $F(x|\lambda)$ normaal, Poisson of binomiaal is, komt het millimeterwerk met empirische verdelingsfunctie een beetje geforceerd over. Beter zou zijn, aandacht te besteden aan methoden om de vorm van $F(x|\lambda)$ te leiden uit symmetrieën in de onvoorwaardelijke verdeling.

Dit boek kan een waardevolle bron zijn voor de specialist in Bayesiaanse methoden, maar niet-specialisten moet het worden afgeraden.

A. Brandt, P. Franken and B. Lisek

Stationary Stochastic Models

John Wiley & Sons, New York, 1990, 344 pag., ISBN 0-471-92132-7, £ 39.95

The contents of this book is centred around recursive stochastic equations of the form

$$X_{n+1} = f(X_n, U_n), \quad n \in \mathbb{Z}. \quad (1)$$

Here the "input" (U_n) , a given stationary sequence of random elements, represents the influence of the environment. The X_n may be interpreted as the consecutive states of a system at some suitably chosen embedded epochs. For many standard and non-standard queueing systems the states X_n , seen by customers arriving at the times T_n , can be described by a model like (1). The state at an arbitrary time t can often be described by a continuous time analogue.

The aim of the book is to present some methods to obtain a (unique) stationary solution of (1) and to derive (as a consequence of these methods) some results for specific stochastic models, especially queueing models. Relationships between stationary sequences and stationary marked point processes (MPP's) offer a very elegant way to relate and to solve a discrete-time stochastic equation (1) and a continuous time analogue.

Throughout this book batch arrivals are allowed. Since it is assumed that the customers arriving at the same time are a-priori ordered, the notion "ordered MPP" is introduced.

Chapter 1 is about existence, uniqueness, and continuity of stationary solutions of the general equation (1). Based on the stationary input (U_n) three methods are discussed to construct stationary, strong or weak solutions. Under certain assumptions one of the methods can be used to construct ergodic weak solutions when starting with ergodic input. Stationary weak solutions are written as a mixture of such ergodic components. Some methods for proving continuity results are discussed.

The case that the input consists of i.i.d. elements is considered in Chapter 2. In an example the existence of non-Markovian solutions of (1) is shown.

In Chapter 3 (ordered) MPP's on the real line are introduced. Since a-priori ordered batch arrivals are allowed, the authors chose to define such MPP's in terms of occurrence-mark couples $[T_n, K_n]$, and not as random counting measures. Special attention is paid to the one-to-one correspondence between a family of distributions of stationary (ordered) MPP's and a family of distributions of stationary sequences $([A_n, K_n])$ with $A_n \geq 0$. (Here A_n should be interpreted as the length of the interval between the n th and $(n+1)$ th occurrence.) The resulting inversion formulae are essential for many parts of the book.

Continuous time models are considered in Chapter 4. If $(X(t))$ represents the state process of a system which changes its states not only at arrivals T_n but also between the arrivals, then it can often be modelled by

$$\begin{cases} X_{n+1} = f(X_n, [A_n, Y_n]), & n \in \mathbb{Z} \\ X(t) = f(X_n, [t - T_n, Y_n]) & \text{if } T_n < t \leq T_{n+1} \end{cases} \quad (2a)$$

$$(2b)$$

It is usually assumed that $f(x, [a, y])$ is left-continuous at a . The input is now repre-

sented by $A_n = T_{n+1} - T_n$ and by other random influence Y_n from outside (e.g. service times). Since $X(T_{n+1}) = X_{n+1}$ if $T_n < T_{n+1}$, the theory of stochastic processes with an embedded (ordered) marked point process (PEMP's) is an appropriate mathematical tool for treating such systems. A stationary weak solution of (2a) can be generalized to a synchronous PEMP which is a solution of both equations (2a, b). This PEMP is usually not stationary. Since a PEMP can be considered as an MPP with special marks, it is however possible (by making use of the one-to-one correspondence mentioned in Chapter 3) to construct a stationary PEMP which satisfies (2a, b).

In Chapter 5 the results of Chapter 1 are applied to some standard and non-standard queues: $G/G/\infty$, $G/G/m/0$, $G/G/1/\infty$, $G/G/m/\infty$ with FCFS, $G/G/m/\infty$ with cyclic queueing discipline, single server queue with warming up, many server loss system with repeated call attempts, and open networks of loss systems. The input $([A_n, Y_n])$ is always stationary and $T_0 \equiv 0$. It is assumed that the mark Y_n comprises some characteristics of the n th customer, such as its required service time, priority number, and so forth. The state X_n seen by the n th customer is described by (2a), where f depends on the queue that is considered. For every queue conditions are given on the input ensuring the existence of a unique stationary solution of (2a). In many cases the state of the system which starts working at $T_0 = 0$ with an arbitrary initial state x evolves to the unique stationary state process.

In Chapter 6 the relationship between arrival-, time-, and departure-stationary queueing processes is considered. It is always assumed that the input $([A_n, Y_n])$ is stationary and ergodic, that $0 < EA_0 < \infty$, and that $f(x, [a, y])$ is left-continuous at a . A stationary weak solution of (2a) (as, e.g., the solutions for the queueing systems in Chapter 5) can be extended to a solution of (2a, b) which can be considered as a synchronous PEMP, the "arrival-stationary model". By the theory of Chapter 4 this leads to a unique stationary PEMP (the "time-stationary model") which satisfies (2a, b). Inversion formulae make it possible to express time-stationary probabilities in terms of arrival-stationary probabilities and vice versa. Some conservation laws, PASTA results, and Little type formulae are derived.

Some special problems concerning batch arrival queues are treated in Chapter 7. "Batch-arrival-stationary" quantities are introduced. The techniques are applied to infinite server and single server queues with batch arrivals.

In Chapter 8 the continuity of queueing models is considered. For the systems $G/G/\infty$, $G/G/m/0$, $G/G/1/\infty$, and $G/G/m/\infty$ conditions are formulated for continuity results for the arrival- and time-stationary solutions of (2a, b).

In Chapter 9 two non-queueing applications of the methods of Chapter 1 are considered. First the stochastic difference equation $X_{n+1} = B_n X_n + C_n$, $n \in \mathbb{Z}$, is investigated. It is assumed that the "input" $([B_n, C_n])$ is stationary and ergodic. Under weak additional assumptions the existence and continuity of the unique stationary solution is shown. Finally, some robust filters for cleaning a stationary sequence from outliers, is discussed. Some mathematical preliminaries about stationary sequences and convergence in distribution are treated in the appendix.

According to the preface the authors consider the well-known book "Queues and Point Processes" by Franken, König, Arndt and Schmidt as a direct predecessor of this book. This is not exaggerated. The books have much in common. Point processes and queueing models

are the main subjects of both books. Notation and organization of the present book clearly originated from the above monograph. There are also differences. Batch arrivals, a-priori ordered, are allowed now and some general methods for treating stationary stochastic models are explained more broadly (Chapters 1-4, 7,9). The notion "Palm distribution", very essential for the above book, is not used anymore. Since a-priori ordered multiple occurrences are now permitted, the Palm approach is replaced by the "stationary sequence approach".

Some points are less satisfactory. As in many books on point processes and queueing theory, notations are very complex. For instance, many symbols are needed to describe the notion PEMP in all its appearances (synchronous, stationary, as a strong solution of (2a, b), as a weak solution of (2a, b)). It should, however, be admitted that it is not easy to avoid this problem. There are many examples in the book, usually about queueing systems. It is, however, never indicated where an example ends and where theory starts again.

The book contains many recent and interesting results on stochastic models, and some new points of view. Although a summary at the beginning of each chapter contains all essential notions and results of that chapter, it takes a lot of time to penetrate into the essence of the book. Still it must be recommended to specialists in applied probability and stochastics, especially those working in point process, queueing, and inventory theory.

Gert Nieuwenhuis

Department of Econometrics, Tilburg University

G.E.P. Box and G.C. Tiao

Bayesian Inference in Statistical Analysis

John Wiley and Sons, New York, 1992 (1973), xviii + 588 pag., ISBN 0-471-57428-7, £ 30.95.

Het boek (een exacte kopie van de oorspronkelijke versie uit 1973) is een nieuwe uitgave in de serie Wiley Classics Library, die als doel heeft klassieke boeken in goedkope versie beschikbaar te houden voor toekomstige generaties wiskundigen en wetenschappers (het is uit deze formulering onduidelijk of Wiley beide groepen disjunct acht). Opname van het boek in deze serie is een waarde-indicatie waarbij ieders individuele mening overbodig wordt, maar laat ik toch opmerken de eer die het boek door deze opname ontvangen heeft volkomen terecht te achten. Twijfels of 'goedkoop' past bij de prijs zijn gerechtvaardigd, maar de omvang en kwaliteit van het werk zorgen dat de koper het boek zeker niet te duur zal vinden. Het boek geeft een zeer uitgebreide en nauwkeurige behandeling van de Bayesiaanse statistiek, ervan uitgaande dat de lezer een goede theoretische kennis heeft van kansrekening, calculus, algebra en frequentistische statistiek. Dit laatste maakt het boek aantrekkelijk voor studie van de Bayesiaanse theorie voor mensen met een 'klassieke' statistiek opleiding (zoals de Nederlandse statistiek opleidingen). Voorts wordt voortdurend geprobeerd de invloed van subjectieve gegevens te minimaliseren, waardoor statistische conclusies vaak overeenkomen met die uit de frequentistische statistiek, maar waarbij

duidelijk wordt dat de Bayesiaanse methode ook uitkomst kan bieden in gevallen waar de frequentistische statistiek faalt. Door deze opzet valt een belangrijk aspect buiten beschouwing, namelijk beslissingstheorie waarbij subjectieve gegevens expliciet gebruikt moeten worden, het terrein waarop de Bayesiaanse methode het duidelijkst grote voordelen biedt. Ook is het jammer dat de bijbehorende opzet, startend bij presentatie en bespreking van de axioma's van subjectieve kansen, ontbreekt. We zouden de aard van het boek pseudo-frequentistisch kunnen noemen, hetgeen door vele lezers met bovengenoemde (en noodzakelijke) achtergrond zeker als een voordeel zal worden beschouwd.

De inleiding (H.1, *Nature of Bayesian Inference*) is uitgebreid, goed leesbaar en geeft een goede indruk van de Bayesiaanse theorie. De verdere indeling van het boek wordt duidelijk door de titels van de hoofdstukken te noemen (steeds worden dezelfde vraagstellingen besproken als bekend uit de frequentistische statistiek):

- H.2 Standard Normal Theory Inference Problems;
- H.3 Bayesian Assessment of Assumptions: 1. Effect of Non-Normality on Inferences about a Population Mean with Generalizations;
- H.4 Bayesian Assessment of Assumptions: 2. Comparison of Variances;
- H.5 Random Effect Models;
- H.6 Analysis of Cross Classification Designs;
- H.7 Inference about Means with Information from more than One Source: One-Way Classification and Block Designs;
- H.8 Some Aspects of Multivariate Analysis;
- H.9 Estimation of Common Regression Coefficients;
- H.10 Transformation of Data.

Intensief wordt gebruik gemaakt van de elegante matrixnotatie, hetgeen een voordeel is met het oog op de jongste generatie software. Natuurlijk is het boek enigszins gedateerd, waardoor belangrijke aspecten als rekentechnieken en robuuste methoden (de belangrijkste ontwikkelingen van de laatste jaren) niet besproken worden. Wanneer het boek gebruikt wordt als eerste kennismaking met de theorie is dit echter geen nadeel, omdat het nodig is eerst de stof uit dit boek te begrijpen alvorens de laatste ontwikkelingen op waarde te kunnen schatten (het lijkt mij hier niet de plaats een overzicht van recente literatuur te geven; geïnteresseerden kunnen zich tot mij wenden).

Samenvattend kunnen we zeggen dat het boek terecht tot klassieker gepromoveerd is. De stof is omvangrijk (voor universitair gebruik denk ik aan cursussen met totale duur van een jaar) en duidelijk gepresenteerd, waardoor het boek geschikt is als cursusmateriaal nadat men de frequentistische statistiek bestudeerd heeft. Ook als naslagwerk is het zeer geschikt, waarbij uit praktische overweging de oorspronkelijke versie (hardback) de voorkeur heeft boven deze paperback uitvoering. Het boek maakt duidelijk dat de Bayesiaanse theorie ten onrechte in Nederland het ondergeschoven kindje van de statistiek is.

F.P.A. Coolen

Faculteit Wiskunde & Informatica, Technische Universiteit Eindhoven

D.M. Eddy, V. Hasselblad, and R. Shachter.

Meta-Analysis by the Confidence Profile Method: The Statistical Synthesis of Evidence

Academic Press, San Diego, 1992, vii + 428 pag., ISBN 0-12-230620-1, \$ 59.95

Inclusief demonstratieversie programma Fast*Pro. Werkversie programma £ 190.50.

Meta-analyse wordt meestal gedefinieerd als een verzameling technieken die gebruikt wordt om de resultaten van het beschikbare empirische onderzoek op een bepaald terrein kwantitatief te analyseren. Het doel is om langs statistische weg een samenvatting of synthese te leveren van de beschikbare onderzoeksresultaten. De statistiek dient daarbij om de gepubliceerde resultaten samen te vatten, en aan te geven in welke mate en onder welke condities de conclusies generaliseerbaar zijn. Meta-analyse verschilt van traditionele essayistische literatuuronderzoek vooral door de systematische en gedetailleerde zoekprocedure en door het gebruik van statistische procedures bij het samenvatten van de gevonden resultaten. Klassieke publikaties zijn Light en Pillemer (1984) over het literatuuronderzoek, en Hedges en Olkin (1985) over de statistische procedures.

Het boek 'Meta-analysis by the confidence profile method' van Eddy, Hasselblad en Shachter staat in een andere traditie. Het belangrijkste doel van de 'Confidence Profile Method' is *niet* het samenvatten van beschikbare resultaten ten behoeve van een literatuuroverzicht. Het doel van de confidence profile method is primair om op basis van de in de literatuur beschikbare informatie *een beslissing te nemen*. Aangenomen wordt dat er sprake is van een probleem waarover een beslissing moet worden genomen, waarbij het uitstellen van de beslissing in afwachting van nader onderzoek geen reële optie is. De confidence profile method is oorspronkelijk ontworpen voor medische toepassingen, en de voorbeelden in het boek betreffen ook alle medische problemen. De methode kan echter ook op andere terreinen worden toegepast: de auteurs noemen zelf landbouwkunde, biologie, ecologie, en sociale wetenschappen.

Het op te lossen probleem wordt in de confidence profile method vertaald in een verzameling parameters. Sommige parameters zijn 'parameters of interest', deze zijn meestal gekoppeld aan uitkomsten-variabelen. Ander parameters zijn 'circumstances of interest', dit kunnen zijn interventies, steekproefkenmerken, designkenmerken, en dergelijke. Uiteindelijk moeten al deze parameters geschat worden. Als hulpmiddel voor de onderzoeker raden de auteurs aan het probleem te beschrijven in de vorm van een 'influence diagram', een padmodel dat aangeeft hoe de verschillende parameters elkaar beïnvloeden. Daarbij wordt onderscheid gemaakt tussen 'basic parameters' (te vergelijken met exogene variabelen), 'functional parameters' (te vergelijken met endogene variabelen), en 'evidence'. De evidence is afkomstig uit de kwantitatieve informatie in de onderzoeksliteratuur: informatie over (verdelingen van) basic en functional parameters.

Het statistische model van de confidence profile method bevat zes elementen: 1) de evidence, 2) de basic parameters, 3) (non-informatieve) priors voor de basic parameters, 4) likelihood functies voor de evidence gegeven de basic parameters, 5) de functional parameters, en 6) transformatiefuncties die de functional parameters beschrijven in termen van de basic parameters en andere functional parameters. Maximum likelihood en Bayesiaanse methoden kunnen gebruikt worden om de parameters en hun standaardfouten te schatten. Het boek beschrijft het algemene model zeer beknopt, en geeft daarna een aantal

speciale gevallen van het algemene model. Interessante speciale gevallen zijn varianten waarin parameters worden ingevoerd om te corrigeren voor veronderstelde onbetrouwbaarheid van metingen of bedreigingen van de interne validiteit van bepaalde onderzoeken. Veel aandacht wordt besteed aan de meta-analyse van (quasi)experimenteel evaluatieonderzoek, waarbij een bepaalde interventie geëvalueerd moet worden. Een beperking van het boek is dat de auteurs vrijwel uitsluitend toepassingen bespreken bij (quasi) experimentele interventiestudies. Ook het computerprogramma Fast*Pro beperkt zich hiertoe.

De confidence profile method is een interessante aanvulling op de meer bekende methoden voor meta-analyse. Net als bij 'gewone' meta-analyse is de verzameling en selectie van de te analyseren literatuur van groot belang. Wanneer de verzamelde literatuur onvolledig is, of methodologisch zwak, of zeer weinig details rapporteert (bekende problemen in de meta-analyse) dan kan de analysemethode daar weinig aan verhelpen. Specifieke problemen bij de confidence profile method zijn het gebruik van subjectieve waarschijnlijkheden bij het specificeren van prior verdelingen. Dit is met name van belang wanneer inderdaad beslissingen genomen moeten worden op basis van (de verdeling van) een uitkomst-parameter. Men kan de invloed van de subjectieve kansschattingen verminderen door een maximum likelihood oplossing kiezen, of een weinig informatieve prior. De auteurs raden aan om hiervoor een gevoeligheidsanalyse uit te voeren. Een aantrekkelijk kenmerk van de confidence profile method, vergeleken met traditionelere vormen van meta-analyse, is dat de onderzoekers gedwongen worden een model op te stellen voor de situatie die zij wensen te beschrijven. De constructie van het 'influence diagram', het padmodel dat de verschillende parameters en hun onderlinge relaties weergeeft, is tegelijkertijd een problematisch punt van de hele analyse, dat in het boek onderbelicht blijft. In bepaalde meta-analyses (bijvoorbeeld analyses waarbij geprobeerd wordt voor potentiële artefacten te controleren) kan het influence diagram behoorlijk complex worden. Het influence diagram is echter niets anders dan een conceptueel model van het onderzochte proces. Negatief omschreven is een conceptueel model niet meer dan een samenhangende verzameling vooroordelen van de onderzoekers. Een gevoeligheidsanalyse voor verschillende influence diagrams zou zeker tot het repertoire moeten behoren.

Het meegeleverde programma Fast*Pro Demo is tamelijk beperkt. Het bevat slechts één rekenmethode, kan geen gegevens van schijf lezen of op schijf opslaan, en kan geen figuren printen. Eenvoudige analyses van de meeste voorbeelden uit het boek kunnen er mee uitgevoerd worden. Onderzoekers die echt met deze methode aan de slag willen, moeten er op rekenen dat ze dan ook de werkversie van het programma zullen moeten aanschaffen.

Referenties

- Hedges, L.V. & Olkin, I. (1985). *Statistical methods for Meta-Analysis*. San Diego: Academic Press.
- Light, R.J. & Pillemer, D.B. (1984). *Summing Up: The Science of Reviewing Research*. Cambridge, MA: Harvard University Press.

J.J. Hox

Faculteit Pedagogische en Onderwijskundige Wetenschappen, UvA

D.C. Montgomery en E.A. Peck

Introduction to linear regression analysis

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

Sinds het standaard werk van Draper en Smith in 1966 verscheen, zijn er tamelijk veel boeken over regressieanalyse geschreven en de vraag kan natuurlijk gesteld worden of dit boek iets bijdraagt aan de reeks van reeds verschenen boeken over lineaire regressieanalyse. Het boek veronderstelt dat studenten beschikken over elementaire statistische kennis, inclusief de meest gangbare significantietoetsen en betrouwbaarheidsintervallen en eenvoudige regels uit de matrix algebra kunnen toepassen. In tegenstelling tot sommige leerboeken over statistische technieken waar de beschrijving van de te gebruiken computerprogramma's in de tekst verweven is met de uitleg van de techniek zelf, wordt in dit boek slechts de uitvoer van SAS en BMDP gepresenteerd als voorbeeld. En zo hoort het dacht ik ook.

Het boek is opgebouwd uit 12 hoofdstukken en 3 appendices. Zoals elk boek over regressieanalyse begint het boek in *hoofdstuk 1 en 2* met een uitleg over enkelvoudige lineaire regressie met hypothese toetsing, betrouwbaarheidsintervallen en voorspelling. Procedures die met de modelfit te maken hebben, zoals residuen-analyse, outlier-detectie, transformaties en gewogen kleinste kwadratenschatting worden in *hoofdstuk 3* behandeld. Vervolgens worden in *hoofdstuk 4* uitgebreid dezelfde onderwerpen voor multiële regressieanalyse besproken. De auteurs geven ook aandacht aan residuen-plots en partiële residuen-plots, PRESS residuen en multicollineariteit.

In *hoofdstuk 5* wordt polynomische regressieanalyse voor een of twee onafhankelijke variabelen besproken en in *hoofdstuk 6* wordt het gebruik van indicatorvariabelen in relatie met categorische variabelen uitgelegd. Dan komen de problemen aan de orde die voor onderzoekers bijzonder belangrijk zijn. In *hoofdstuk 7* wordt het probleem van selectie van variabelen in multiële regressieanalyse uitgebreid en duidelijk besproken en voorbeelden worden aan de hand van BMDP programma's toegelicht. Elke onderzoeker die multiële regressieanalyse heeft uitgevoerd is wel eens tegen het probleem van multicollineariteit opgelopen. In *hoofdstuk 8* wordt aandacht gegeven aan methoden om multicollineariteit vast te stellen en om ermee om te kunnen gaan, zoals ridge regressie-, principale componenten regressie- en eigenwaarden regressieanalyse. *Hoofdstuk 9* behandelt autocorrelatie en geeft uitleg over gegeneraliseerde en gewogen kleinste kwadratenschattingsprocedures. Ook worden robuuste schattingsprocedures (M, R en L schatters) besproken. Hoewel deze technieken duidelijk worden uitgelegd, had ik persoonlijk graag gezien dat hier meer over verteld werd. De relatie tussen regressieanalyse en variantieanalyse, welke van belang is bij ongebalanceerde experimentele designs, wordt besproken. Ook een wat uitgebreidere behandeling van het optimale ontwerpen van experimenten voor regressieanalyse zou hier mijns inziens gewenst zijn. Tenslotte wordt in hetzelfde hoofdstuk kort ingegaan op niet-lineaire regressieanalyse en is *hoofdstuk 10* gewijd aan validatie van het regressiemodel met bijzondere aandacht voor het DUPLEX algoritme.

Bij elk hoofdstuk zijn een aantal oefenopgaven gegeven, die meestal samen met de in de tekst uitgewerkte voorbeelden en het gebruik van SAS of BMDP gemaakt kunnen

worden. In een appendix worden ook de datasets die bij de oefenopgaven gebruiken kunnen worden gegeven. De voorbeelden worden met redelijk veel grafieken in de tekst uitgewerkt. In de tekst worden nauwelijks afleidingen gegeven en afgezien misschien van het gebruik van de elementaire matrix algebra notatie is het boek goed te gebruiken voor studenten, AIO's en onderzoekers uit de sociale wetenschappen, economie, bedrijfskunde of een technische richting, die primair belangstelling hebben voor het correct gebruiken van regressieanalyse procedures.

Dit boek heeft door de lay-out, de eenvoudige en duidelijk notatie en de kort en bondige uitleg van de stof een haast klassieke stijl. Terugkomend op de vraag of dit boek iets bijdraagt aan de reeds verschenen werken op het gebied van regressieanalyse, is er natuurlijk vast te stellen dat de behandelde stof in vergelijkbare boeken ook te vinden is. Persoonlijk vind ik echter dat dit boek aantrekkelijk is vanwege de bondige en duidelijk uitleg die gecombineerd wordt met uitvoer uit de computerprogramma's SAS en BMDP. Dit maakt het boek zeer geschikt als tekstboek voor een inleidende cursus over regressie-analyse.

Referentie

Draper, N.R. and H. Smith (1966, 1981) *Applied Regression Analysis*, Wiley, N.Y.

M.P.F. Berger

Universiteit Twente

D.C. Hoaglin, F. Mosteller & J.W. Tukey

Fundamentals of exploratory analysis of variance

John Wiley & Sons, New York, 1991, xvii + 430 pag., ISBN 0-471-52735-1, £ 43.95

Het boek is opgebouwd uit verschillende bijdragen van in totaal 11 auteurs. Het centrale thema is het analyseren van data, welke wij gewoonlijk doen met een variantie analyse tabel. De auteurs presenteren echter een techniek welke meer aansluit bij het EDA concept van Tukey, door middels tabellen de opbouw van de waarde van de gegevens uiteen te rafelen in de bijdrage van de diverse componenten. Daarmee worden de getallen waarop de samenvattende kentallen in de ANOVA tabel zijn gebaseerd, zichtbaar gemaakt. Deze techniek is zeker niet nieuw. In de meeste praktijkgerichte boeken over proefopzetten staan wel sommen van matrices met het algemeen gemiddelden, het gemiddelde per behandeling en de residuen. Maar vaak zijn dit voorbeelden van een één-factor proefopzet of hoogstens een twee-factor schema en wordt in de rest van het boek alleen nog gesproken over de variantie analyse tabel.

Om de methode te illustreren is ruim gebruik gemaakt van voorbeeldgegevens uit de dagelijkse praktijk, wat de praktische waarde van dit boek aanmerkelijk vergroot.

In het eerste hoofdstuk over de algemene principes van variantie analyse (**Concepts and examples in analysis of variance**, J.W. Tukey, F. Mosteller & D.C. Hoaglin), wordt al direct

gewerkt met kleine subtabellen (in de tekst 'overlays' genoemd) waarin de uiteenrafeling van de gegevens naar factor wordt gepresenteerd. Aan de hand daarvan wordt ook het interactie concept geïntroduceerd. Tevens worden hier de eerste voorbeelden uit de doeken gedaan.

Het tweede hoofdstuk (**Purpose of analyzing data from the analysis of variance**, *F. Mosteller & J.W. Tukey*) wordt ingegaan op het waarom van analyseren van data. Het begrijpen van het onderliggende fenomeen, de aard van de variabelen en de invloed van diverse factoren. De geboden structuren en gereedschappen moeten leiden tot een beter begrip van het fenomeen dat bestudeerd wordt. Het is een hoofdstuk met een algemene strekking.

In **Preliminary examination of data** (*F. Mosteller & D.C. Hoaglin*) wordt een korte verhandeling gegeven over eenvoudige EDA technieken. Grafieken, box-and-whisker plots en stam-en-blad diagrammen komen hier aan de orde. Doel is om vooraf inzicht te krijgen in de data, onmogelijke waarden te achterhalen en mogelijke uitschieters te identificeren voordat de verdere analyse wordt uitgevoerd.

Na hoofdstuk 4 (**Types of factors and their structural layout**, *J.D. Singer*) waarin onderscheid gemaakt wordt tussen geordende en ongeordende variabelen en gemeten variabelen, gekruiste en geneste factoren, wordt in hoofdstuk 5 en 6 (**Value splitting, taking the data apart**, *C.H. Schmid* en **Value splitting involving more factors**, *K.T. Halvorsen*) een deel van de feitelijke boodschap gepresenteerd. De tabel met de responsies wordt in deze hoofdstukken ontleed in diverse subtabellen (de 'overlays'), gekoppeld aan de niveaus van de factoren, waar de bijdrage van deze factoren staat vermeld.

Bij meer factoren wordt het geheel niet eenvoudiger om te presenteren, maar door een geschikte rangschikking van de subtabellen wordt toch inzicht verkregen. Het eenvoudigst is het presenteren van deze techniek bij geneste factoren. Dan is de analyse van een genest schema, bijvoorbeeld bij een onderzoek naar meetmethode onnauwkeurigheid, aan niet statistisch geschoolde medewerkers goed duidelijk te maken door de aanschouwelijkheid van de methode.

Hoofdstuk 7 (**Mean squares, F tests and estimates of variance**, *F. Mosteller, A. Parunak en J.W. Tukey*), hoofdstuk 8 (**Graphical display as an aid to analysis**, *J.D. Emerson*) en hoofdstuk 9 (**Components of variance**, *C. Brown en F. Mosteller*) geven een overzicht van de analyse methoden zoals deze in veel boeken voorkomt. De grafische analyse presenteert effecten plaatjes, waarbij de effecten van alle factoren, dus ook de residuen, in een enkele grafiek worden weergegeven. Daarbij wordt aandacht besteed aan het feit dat de geschatte effecten niet allemaal dezelfde variantie hebben, maar wel in één plaatje moeten worden gepresenteerd om de overzichtelijkheid te behouden. In hoofdstuk 9 worden niet alleen de variantie componenten geschat, maar worden ook (verschillende) betrouwbaarheidsintervallen gepresenteerd.

Hoofdstuk 10 (**Which denominator**, *T. Balckwell, C. Brown en F. Mosteller*) handelt over het toestingsschema in de variantie analyse tabel, afhankelijk van de aard van de factoren (vast of random) en de vorm van de schema's (geneste factoren of gekruiste factoren). Ook wordt aandacht besteed aan het vormen van een geschikte noemer voor de F-test als som van gemiddelde kwadraten uit de variantie analyse tabel als deze niet direct in de tabel

voorhanden is.

In hoofdstuk 11 (**Assessing changes**, *J.W. Tukey, F. Mosteller en C. Youtz*) wordt het bijelkaar vegen van factoren, omdat de invloed niet te onderscheiden is van ruis, behandeld. Dat bij elkaar vegen van factoren gaat niet zoals 'gebruikelijk' van hoogste orde interactie naar hoofdeffect, maar in de hier gepresenteerde methode van hoofdeffect naar hoogste orde interacties, dus in de andere richting. Het criterium is de ' $F > 2$ regel' op basis waarvan effecten bij elkaar worden geveegd als de F-waarde kleiner is dan 2. In een twee factor schema met interactie komt dat dan neer op het bijelkaar vegen van twee-factor-interactie en residuen als het gemiddeld kwadraat voor de interactie kleiner is dan 2 maal het gemiddeld kwadraat van de residuen. Daaropvolgend wordt dan gekeken of het gemiddelde kwadraat van een hoofdeffect (A of B) kleiner is dan 2 maal het gemiddelde kwadraat van het (dan ontstane) residu, en zo voort.

Als commentaar op deze regel geven de auteurs zelf aan dat het bij elkaar vegen van effecten niet betekent dat het betreffende effect niet bestaat, maar dat het onaantrekkelijk is om apart te laten bestaan gezien het geringe verschil met ruis.

Deze aanpak doet nogal vreemd aan als je bent grootgebracht met p-waarden, maar open staan voor deze aanpak en er enige tijd over nadenken kan leiden tot frisse ideeën. Zo zaligmakend zijn die p-waarden nu ook weer niet.

In hoofdstuk 12 (**Quantitative and qualitative confidence**, *J.W. Tukey en D.C. Hoaglin*) wordt gesproken over de verschillen tussen de niveaus van de factoren en hoe je die kunt achterhalen. De standaardprocedures als Bonferroni en Studentized range komen daar aan bod, overigens nu wel weer gebaseerd op nette onbetrouwbaarheidsdrempels.

Tenslotte wordt in hoofdstuk 13 (**Introduction to transformations**, *J.D. Emerson*) een motivatie en enkele technieken gegeven voor het transformeren van de data naar een andere schaal.

De eigenheid van het boek is dus hoofdzakelijk gegroepeerd in hoofdstuk 5, 6, 11 en 12. De andere hoofdstukken geven overzichten van de 'traditionele' aanpak bij variantie analyse. Het boek kan toch nieuwe inzichten leveren en daarom is het volgens mij goed dat een ieder die veel van de variantie analyse techniek gebruik maakt en dat vooral ook doet voor niet statistisch onderlegde klanten, kennis neemt van de aangeboden methoden. Zoals reeds gesteld verschaft de methode van ontbinding van de gegevens in 'overlays' inzicht in wat nu precies in zo'n (samenvattende) variantie analyse tabel wordt weergegeven. Dat kan voor de klanten motiverend werken bij een volgende toepassing.

Als leerboek voor de techniek van variantie analyse vind ik het boek minder geschikt. Het is ook niet als zodanig opgebouwd en de presentatie van deze alternatieve methode voor het analyseren van variantie analyse data gaat ervan uit dat de standaard variantie analyse techniek voldoende beheerst wordt. Het boek is wel geschikt om eens op een andere wijze naar de analyse methode te kijken.

Arend Oosterhoorn

H. Leon Harter

The chronological annotated bibliography of order statistics, vol VI: 1966-1967

American Science Press, Syracuse, 1991, vi + 232, ISBN 0-935950-24-9, \$ 95.00

H. Leon Harter

The chronological annotated bibliography of order statistics, vol VII: 1968-1969

American Science Press, Syracuse, 1991, vi + 292, ISBN 0-935950-25-7, \$ 95.00

Reeds in KM 40 schreef ik u over de eerste vijf delen, waarvan deel III, IV en V in die aflevering van Kwantitatieve Methoden zijn besproken. Nu zijn de volgende twee delen ook uitgebracht. Naast 308 publicaties uit 1966 en 356 publicaties uit 1967 bevat deel VI ook nog 71 bijdragen uit de periode voor 1966.

Het aantal publicaties uit 1966 en 1967 wordt aangevuld in deel VII waarna een groot aantal uit de jaren 1968 en 1969 wordt besproken.

Voor verdere uitwijding over de werkwijze en de waarde van deze delen verwijs ik naar de bespreking in KM40.

Arend Oosterhoorn

>>>>>>>>>> **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.

Boeken waarvoor een $\boxtimes$ symbool staat zijn nog niet 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.

$\boxtimes$ **Todorovic, P.**

An introduction to stochastic processes and their applications

Springer Verlag, Berlijn, 1992, xiv + 289 pag., ISBN 3-540-97783-X, DM 98.00

Bachem, A. & W. Kern

Linear Programming Duality, an introduction to oriented matroid

Springer Verlag, Berlijn, 1992, iv + 216 pag., ISBN 3-540-55417-3, DM 68.00

Lindvall, T.

Lectures on the coupling method

John Wiley & Sons, New York, 1992, xii + 257 pag., ISBN 0-471-54025-0, £ 39.95

- ♣ Bofinger, E., Dudewicz, E.J., Lewis, G.L. & Mergens, K.
The frontiers of modern statistical inference procedures, II
 American Sciences Press, Syracuse, 1992, xii + 498 pag., ISBN 0-935950-30-3, \$ 98.75
- ♣ Arnold, B.C., N. Balakrishnan & H.N. Nagajara
A first course in order statistics
 John Wiley & Sons, New York, 1992, xvii + 279 pag., ISBN 0-471-57416-3, £ 39.95
- ♣ Scott, D.W.
Multivariate density estimation, theory, practice and visualization
 John Wiley & Sons, New York, 1992, xii + 317 pag., ISBN 0-471-54770-0, £ 47.50
- Schreuder, F.
SPSS/PC+ tables
 Academic Service, Schoonhoven, 1992, 122 pag., ISBN 90-5261-086-X, f 35.00
- ♣ Kleijnen, J. & W. van Groenendaal
Simulation, a statistical perspective
 John Wiley & Sons, New York, 1992, x + 241 pag., ISBN 0-471-93055-5, £ 16.50

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

Vakantiecursus 1992: Systeemtheorie

Stichting Mathematisch Centrum, Amsterdam, 1992, iii + 209 pag., ISBN 90-6196-409-1,
 f 60.00.

CWI syllabus 31 is de neerslag van de vakantiecursus 1992 met als onderwerp de systeem-
 theorie. De bundel presenteert de bijdragen van de diverse docenten tijdens deze cursus.

Introduction to Mathematical System Theory, *G.J. Olsder*

Hulpmiddelen uit de Lineaire Algebra, *A.W. Grootendorst*

Inleiding Gewone Differentiaalvergelijkingen, *H.J.C. Huijberts*

Stochastiek, *J.Th.M. Wijnen*

Sturen en Waarnemen, *J.W. van der Woude*

Tijdoptimale Besturing van Lineaire Systemen, *M.L.J. Hautus*

Kalman Filtering, *A.W. Heemink*

Recent Developments in Mathematical System Theory, *G.J. Olsder*

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 nauw-
keurig omschrijven)
