

BOEKBESPREKINGEN

Redactie:

J.J. Dik

Correspondentie-adres:

Volvo Car bv
Postbus 1015
5700 MC Helmond

Telefoon: (04920) 31368

Constance Reid

Neyman from life

Springer Verlag, New York, 1982, \$ 20.-

Dit is een kleurrijk boek over een kleurrijk mens en iedereen met belangstelling voor de mathematische statistiek en zijn geschiedenis moet het beslist lezen. Het is een eerbewijs aan een groot geleerde en de stichter van het beroemde Statistics Department in Berkeley, een man met principes en een tomeloze werklust, die veel van anderen eiste maar hen dan ook steeds hielp, die strijdbaar was maar altijd hoffelijk, de wereld ernstig nam maar met een knipoog. Na haar succesvolle biografieën van Hilbert en Courant, heeft Mrs Reid ditmaal een levende legende voor de lens gezet. Tijdens de vele gesprekken die zij de laatste jaren van zijn leven met hem voerde, heeft zij hem goed leren kennen en dat is in haar boek voortdurend te merken. Dit is een portret van Jerzy Neyman ten voeten uit.

Hoewel Neyman in 1894 in Rusland werd geboren en Polen in die tijd niet als een zelfstandige staat bestond, heeft hij zichzelf altijd als een geboren Pool beschouwd. In 1912 ging hij in Charkov studeren, waar al snel duidelijk werd dat hij geen groot talent voor de experimentele natuurkunde bezat; apparatuur placht in zijn handen uiteen te vallen. Mrs Reid vertelt een mooi verhaal over zijn onkunde om nietjes in een nietmachine te doen: "I used to think that maybe just American-born could do it, but then somebody suggested that only Poles couldn't do it and then Mark Eudy suggested maybe only one Pole couldn't do it". Zoals wel meer mislukte natuurkunde studenten, richtte hij zich op de wiskunde en zijn meest geliefkoosde onderwerp werd de theorie van de Lebesgue integraal. Als student won hij een prijs voor een verhandeling op dit terrein en nadat hij in 1921 naar Warschau was vertrokken, werd het één van de onderwerpen waarop hij daar in 1924 promoveerde. In de tussenliggende jaren had hij de wereldoorlog en de revolutie meegemaakt, waarbij hij herhaaldelijk in de gevangenis beland was. Vanaf 1921 werkte hij bij het Nationale Landbouwkundige Instituut, waar hij zichzelf statistiek begon te leren. In 1925 kreeg hij een beurs van de Poolse regering om een jaar bij Karl Pearson te studeren aan University College in Londen.

Dit eerste bezoek aan Engeland werd niet bepaald een groot succes. Men vond hem kennelijk maar een vreemde figuur. Karl Pearson was er door de voortdurende ruzies met Fisher niet toegankelijker op geworden: "This may be true in Poland, Mr. Neyman, but it is not true here!" was zijn reactie toen Neyman trachtte uit te leggen dat "onafhankelijk" niet hetzelfde is als "ongecorreleerd". Op zijn beurt vond Neyman bij een bezoek aan Parijs, Borel en Lebesgue heel wat interessanter dan Karl Pearson. Er was eigenlijk maar één reden waarom hij niet tot zijn oude liefde, de maattheorie, terugkeerde: hij was bevriend geraakt met Egon Pearson en hun samenwerking zou hem definitief voor de statistiek winnen.

De jaren die volgden, tot 1934 in Polen en daarna tot 1938 bij Egon Pearson in Londen, zijn in wetenschappelijk opzicht de meest vruchtbare en opwindende van Neyman's carrière. Hij bouwt aan de fundamenten van de mathematische statistiek: de Neyman-Pearson theorie voor het toetsen van hypothesen, betrouwbaarheidsintervallen, stratificatie, randomisering, etc.. Door eminente statistici zoals Lehmann en LeCam aan het woord te laten, weet Mrs Reid de grote betekenis van dit werk duidelijk te maken. Citaten uit de brieven van Neyman aan Egon Pearson, die door Pearson gelukkig zijn bewaard, maken duidelijk hoe moeilijk het was de juiste wiskundige formulering voor de statistische problemen te vinden. Het werk over betrouwbaarheidsintervallen leidde tot een betreurenswaardige controverse met Fisher, met wie Neyman aanvankelijk goede relaties onderhield.

In 1938 vertrekt hij definitief naar Berkeley. Weer komt er een wereldoorlog tussenbeide en pas daarna kan hij zich in alle ernst gaan wijden aan de opbouw van het Statistics Department en het Statistical Laboratory. Hij weet de beste mensen aan te trekken, bezoekers en studenten uit de hele wereld stromen toe en de Berkeley Symposia worden unieke bijeenkomsten waar vertegenwoordigers van vele scholen in de statistiek elkaar ontmoeten. Zijn eigen belangstelling gaat meer en meer uit naar de toepassingen van de statistiek. Als hij in 1961 tenslotte emeritus wordt, verandert er eigenlijk niets. Ieder jaar opnieuw stelt de universiteit hem aan als "Professor Recalled to Active Duty and Director of the Statistical Laboratory" en tot zijn dood in 1981 blijft hij actief. Ook na zijn dood gaat het Neyman Seminar gewoon door in dezelfde zaal. Op de muur staat nu zijn favoriete uitspraak "Life is complicated, but not uninteresting".

W.R. van Zwet, R.U.Leiden

Michael J. Greenacre

Theory and Applications of Correspondence Analysis

Academic Press, London, 1984, 364p, ISBN 0-12-299050-1, £ 34,-

Op het gebied van de correspondentie analyse (CA) en aanverwante technieken bestaat er een groot gat in de markt. Strikt genomen is de wiskunde van CA al in de dertiger jaren afgeleid, maar pas in de jaren zestig is CA voor het eerst tot enige bloei gekomen. Helaas gebeurde dit in Frankrijk, hetgeen tot gevolg had dat er weinig verspreiding plaatsvond. Dit omdat Fransen er haast niet toe te krijgen zijn in het Engels te schrijven en ze ook in het Frans hun werk weliswaar fraai, maar weinig toegankelijk publiceren. Het standaardwerk van Benzecri over CA is hier een goed voorbeeld van. In Nederland heeft met name de afdeling datatheorie van de RUL bijgedragen tot CA en verwante technieken, maar hun chef d'oeuvre, gepubliceerd onder de schuilnaam Albert Gifi, is ook niet direct een goed leerboek. Het boek van Greenacre pretendeert dit wel te zijn en het zou het werk van de Franse school voor de rest van de wereld moeten ontsluiten. Het bevat de volgende hoofdstukken.

1. Introduction
2. Geometric Concepts in Multidimensional Space
3. Simple Illustrations of Correspondence Analysis
4. Theory of Correspondence Analysis and Equivalent Approaches
5. Multiple Correspondence Analysis
6. Correspondence Analysis of Ratings and Preferences
7. Use of Correspondence Analysis in Discriminant Analysis, Classification, Regression and Cluster Analysis
8. Special Topics
9. Applications of Correspondence Analysis

Strikt genomen is slechts vrij weinig voorkennis nodig op het gebied van analyse en lineaire algebra. In hoofdstuk 2 worden alle wiskundige begrippen behandeld die bij CA van toepassing zijn. Hierop wordt in latere hoofdstukken vaak teruggegrepen. Daarom kan ook de geschoolde lezer dit hoofdstuk niet overslaan. In hoofdstuk 3 wordt uitvoerig een voorbeeld behandeld dat in hoofdstuk 4 en 5 als rode draad fungeert. Jammer is daarbij dat de meest fundamentele tabel (3.1.) misleidend is weergegeven in de zin dat de kolom-categoriën (1) t/m (4) corresponderen met (a) t/m (d) in de tekst.

Terecht wordt als basisprincipe uitgegaan van de -metriek, omdat hieruit het beste blijkt wanneer CA toepasbaar is: daar waar de -afstand een zinvolle ongelijkheidsmaat is. In de hoofdstukken 4 en 5 put de auteur zich uit zo helder en schematisch mogelijk de vele theoretische invalshoeken van CA weer te geven. Hij slaagt daar m.i. redelijk in, al heb ik twee dingen gemist. In de eerste plaats geeft hij niet de zeer eenvoudige benadering van CA als het maximaliseren van een correlatiecoëfficiënt; dit staat ergens weggemoffeld bij canonieke correlaties. In de tweede plaats geeft hij geen uitwerking van de -afstand bij multiële CA, iets wat zeer verhelderend in termen van verzamelingen kan worden uitgedrukt. Hoofdstuk 6 gaat over het verdubbelen van tabellen en de relatie met afstand-georiënteerde schaaltechnieken. Dit laatste onderwerp komt matig uit de verf, omdat niet diep genoeg op de materie kan worden ingegaan. Hoofdstuk 7 laat CA zien als tussenstap bij andere technieken. In hoofdstuk 8 worden enkele losse onderwerpen behandeld. Dit gebeurt tamelijk uit de losse pols, hetgeen soms onbevredigend is. Zo worden bij de jackknife niet de basisformules van deze techniek gegeven. Bij het hoefijzer-effect blijft de achterliggende tweedimensionale normale verdeling onbesproken, evenals de eigenschappen van totaal positieve tabellen. In hoofdstuk 9 wordt een diversiteit aan onderwerpen besproken, die een goed beeld geeft van de kracht en de veelzijdigheid van CA.

CA kan men niet alleen uit een boek leren. Het is bij uitstek een techniek waarvan men pas in de praktijk de waarde leert kennen. Dit boek zou echter een uitstekende rol kunnen vervullen als leerboek naast een praktikum en als naslagwerk voor onderzoekers. Ondanks de bovengenoemde omissies is de auteur er zeer behoorlijk in geslaagd toegankelijkheid te combineren met een redelijke theoretische grondigheid. Het gat in de Angelsaksische markt lijkt met dit boek dan ook voorlopig gevuld.

D. Sikkels, CBS

Helmuth Späth

Cluster dissection and analysis, theory, FORTRAN programs, examples

Ellis Horwood Ltd., Chichester, 1985, 226p, ISBN 0-85312-736-0, £ 25,-
(Vertaald uit het Duits: Cluster-Formation-und-Analyse, 1983)

Het boek heeft drie delen: I een wiskundige onderbouwing van methoden voor clusteranalyse, II algoritmen en hun implementatie door middel van FORTRAN IV subroutines en III een aantal hoofdprogramma's met enige toepassingen ervan. De partitie van elementen in clusters wordt gebaseerd op een doel-functie, waarbij aan elke partitie een niet-negatief reëel getal wordt toegekend waardoor een vergelijking tussen de verschillende partities mogelijk is. Niet besproken worden bijvoorbeeld methoden die gebaseerd zijn op hiërarchische methoden, dichtheidsschattingen en ordening van afstanden. De auteur neemt nadrukkelijk afstand van standaard-pakketten en de daarin gebruikelijk toegepaste methodieken. De naam CLUSTAN wordt niet genoemd, evenmin de ontwerper van dit pakket D. Wishart! In het theoretische gedeelte worden wiskundige afleidingen gegeven voor

optimaliteits-eigenschappen voor onder andere het minimum variantie criterium, het minimum determinant criterium, het L_1 criterium voor kwantitatieve, binaire en ordinale data en clusterwijze lineaire regressie. Eigenschappen van bijvoorbeeld minimale afstand methode, het uitwisselen van elementen tussen clusters en de invloed van schaling van de variabelen worden afgeleid bij het hanteren van de diverse criteria.

In het tweede deel van het boek worden listings van een twintigtal subroutines gegeven en besproken. Een aantal hoofdprogramma's die deze subroutines gebruiken worden gepresenteerd in het derde deel. Deze worden vervolgens op geconstrueerde - en tamelijk gekunstelde - data toegepast. De resultaten worden weergegeven in de vorm van uitgebreide - en slecht leesbare - tabellen van getallen en duidelijke twee-dimensionale scatter-plots. De -bekende- afhankelijkheid van het aantal clusters wordt hierbij zichtbaar gemaakt. De auteur trekt bijvoorbeeld de conclusie dat het variantie criterium tendeeert bolvormige clusters te vinden of te vormen, terwijl het determinant criterium vooral ellipsvormen vindt. Hoofdprogramma's, subroutines en test data zijn verkrijgbaar op een magneetband; er wordt niet vermeld tegen welke prijs. De programma's zijn door de auteur gedraaid op een TR440, met vermelding van gebruikte rekentijden. Er worden aanwijzingen gegeven voor implementatie op een microcomputer (microsoft FORTRAN-80).

Voor wie is het boek bestemd? De auteur zegt zelf dat het niet alleen voor wiskundigen bestemd is. Ik heb daar enige twijfel over. Niet alleen is de stijl typisch wiskundig, stel een criterium, geef stellingen en bewijs ze precies, maar vooral lijkt het me dat het gros van de toepassers van clusteranalyse weinig aansluiting vindt bij de methoden die ze gewend zijn toe te passen en de wijze waarop deze worden gepresenteerd in het CLUSTAN manual, het boekje Cluster Analysis van Everitt of het rapport Cluster-Analyse van de (voormalige) Contactgroep Statistische Programmatuur.

Een positieve kant van het boek is wel dat de toepasser wordt geconfronteerd met een aantal methodieken die hij niet in de standaard pakketten tegenkomt. Bovendien verdiept de strakke afleiding van eigenschappen van de methoden het inzicht en haalt clusteranalyse wat uit de sfeer van willekeur in groeperingsmethoden en gelijkens/afstandsmaten. Het presenteren van de listings van subroutines en hoofdprogramma's lijkt me alleen zinvol voor diegenen die precies willen weten hoe het achterliggende algoritme in elkaar zit. Het doorwerken van een FORTRAN-programma zal echter weinigen plezier verschaffen. Voor echte geïnteresseerden lijkt het efficiënter - en plezieriger - de magneetband met programma's aan te schaffen en toe te passen op de gegevens of eigen data om de diverse methoden te vergelijken met elkaar en/of met in CLUSTAN opgenomen methoden.

Concluderend: een boek vooral voor de fijnproever en mogelijk voor de toepasser van clusteranalyse ter confrontatie met een theoretisch fundament.

L.Th. van der Weele, RUG

Vic Barnett & Toby Lewis

Outliers in Statistical Data

John Wiley & Sons, New York, 1984, 2nd ed, xiv+463p, ISBN 0-471-90507-0, £ 25,95

Dit is de tweede druk van een zeer geslaagd boek over een interessant onderwerp. Het biedt een vrijwel volledig overzicht over de literatuur op het gebied van de statistische aanpak van uitschieters. Het is zo geschreven dat de praktische gebruiker de besproken methoden ook direct kan toepassen. Daartoe zijn niet minder dan 37 tabellen opgenomen.

Voor de eerste druk, verschenen in 1978, vormden zo'n 400 artikelen de basis. In de afgelopen zes jaar is de belangstelling voor uitschieters weer

sterk toegenomen; er zijn in die tijd meer dan 300 nieuwe artikelen gepubliceerd. Dit heeft er toe geleid dat de omvang van het boek met 100 pagina's is toegenomen. Er zijn twee nieuwe hoofdstukken toegevoegd over uitschieters in tijdreeksen en over uitschieters in die situaties waar de waarnemingen richtingen zijn in het platte vlak en in de ruimte. Verder is meer uitgebreid aandacht besteed aan robuuste methoden in het geval van de mogelijke aanwezigheid van uitschieters en aan informele, vaak grafische methoden. Bij regressieproblemen wordt er de laatste jaren erg veel gedaan aan het onderwerp "influential observations": waarnemingen die door de waarden van de regressoren een bijzonder grote invloed hebben op de parameterschattingen. In het hoofdstuk over regressie en het lineaire model wordt het verband met uitschieters behandeld. Zeer bewonderenswaardig tenslotte is de wijze waarop de nieuwste literatuur door het hele boek heen zo is verwerkt dat het weer volkomen up to date is geworden.

Zowel voor de practiserende statisticus, die in zijn werk met uitschieters te maken krijgt (en voor wie is dat niet het geval?), als voor diegenen die op dit terrein onderzoek verrichten, kan ik dit boek zonder voorbehoud aanbevelen.

R. Doornbos, THE

A. Buijs

Statistiek om mee te werken

Stenfert Kroese, Leiden, 1984, 324p, ISBN 90-207-1295-0, f43,45

Het hier besproken boek is een tweede, herzien druk, waarbij door de auteur gebruik is gemaakt van een groot aantal reacties die hem bereikt hebben van de zijde van de docenten die de eerste druk gebruikten bij hun cursussen. De auteur schrijft dat het boek bedoeld is als een kennismaking in de breedte, waardoor de student in staat wordt gesteld om het veld enigszins te overzien; voorts vermeldt de auteur dat getracht is om de wiskundige achtergrond van de methodes niet te verwaarlozen, maar dat anderzijds concessies zijn gedaan om de stof redelijk toegankelijk te laten zijn voor studenten die geen wiskunde in het eindexamenpakket VWO hebben. Al met al wordt ernaar gestreefd de onderwerpen aan de orde te laten komen die doorgaans worden behandeld in een cursus elementaire statistiek, meer in het bijzonder voor economen.

De indeling van de stof is als volgt: na een tweetal hoofdstukken over het ordenen, presenteren en samenvatten van gegevens, komt in de volgende vijf hoofdstukken de kansrekening en de elementaire verdelingstheorie aan de orde. Vervolgens worden de onderwerpen toetsen en schatten, zonnodig met gebruikmaking van de t-verdeling, behandeld. De eerste tien hoofdstukken worden afgesloten met een aantal facetten van regressie en correlatie. Het boek bevat dan tenslotte nog een drietal minder algemene hoofdstukken over indexcijfers, tijdreeksen en toepassingsmogelijkheden van de chi-kwadraattoets.

Het boek is geschreven in een plezierige, soepele stijl. Er worden veel voorbeelden gegeven, en er zijn weinig storende drukfouten, hetgeen het boek als studieboek heel geschikt maakt voor de gebruikers. Wel zijn de eerste acht hoofdstukken duidelijk evenwichtiger dan de laatste vijf, waar de noodzaak om veel techniek te introduceren soms wat abrupte stappen oplevert en soms wat nodeloze herhalingen. Ook kan men wat redekavelen over de wiskundige "diepgang"; met name de integraalrekening zal sommige studenten rauw op de maag vallen.

In het volgende worden nog een aantal kritische kanttekeningen geplaatst, waarin zeker enige elementen van persoonlijke voorkeur schuilen.

Voor wat betreft de basishoofdstukken valt het op dat in het geheel geen aandacht wordt geschonken aan het analyseren van tabellen. Ik denk hierbij bijvoorbeeld aan het terugrekenen van gewichten bij tabellen met percentages voor deelgroepen en een totaal. Een begrip als statistische samenhang bij nominale variabelen komt niet expliciet aan de orde; ook zonder associatiematen te introduceren valt hier nog wel iets zinnigs over te zeggen.

De auteur streeft duidelijk naar het systematisch gebruik van bepaalde letters en symbolen met een vaste betekenis, hetgeen in een elementair leerboek alleen maar is toe te juichen. De hoofdletter P wordt gereserveerd voor een kans. Daarom is het nogal ongelukkig om op blz. 106 de notatie $P_{ij} = P(k = k_i, m = m_j)$ te hanteren in plaats van P_{ij} . Voor de parameter van de alternatieve verdeling wordt later in het boek de letter π gehanteerd, maar voordat het begrip steekproef fractie is geïntroduceerd ook wel p. Dit is extra jammer omdat het verschil tussen de parameter en die steekproef fractie, zelfs met de onderstreping van een stochastische variabele, voor studenten vrij lastig is.

Een enkele keer schiet de soepele stijl iets te ver door. Zo wordt covariantie omschreven als (de aanwezigheid van) een zekere samenhang. Het woordje zekere is hier voor tweeërlei uitleg vatbaar: een specifieke vorm van samenhang (i.c. lineaire) óf het complement van onafhankelijkheid hetgeen juist niet bedoeld is. De zaak blijft voor studenten duister omdat wel wordt aangetoond dat onafhankelijkheid covariantie "nul" impliceert, maar niet wordt ingegaan op het feit dat deze implicatie niet omkeerbaar is. Op blz. 114 wordt gesteld dat de formule van de kansdichtheid van de normale verdeling dusdanig ingewikkeld is, dat het niet eenvoudig is de integraal ervan te berekenen. De implicatie "ingewikkeld dus niet integreerbaar" dringt zich op. Op blz. 117 verschijnt het begrip steekproef gemiddelde, maar eerst op blz. 167 wordt het begrip steekproef gedefinieerd; misschien een iets te groot beroep op de intuïtieve betekenis van deze termen.

Een beetje overbodig, wellicht, is de introductie van zowel de keuringskarakteristiek als de grafiek van het onderscheidingsvermogen van een toets, of in ieder geval verwarrend zonder toelichting op hun onderlinge relatie.

Bij de behandeling van de toets voor het verschil van twee gemiddelden op blz. 211 e.v. wordt de F-toets wel erg abrupt geïntroduceerd, zeker in verhouding tot de behoedzame wijze waarop dat in een eerder gedeelte met de t-toets is gebeurd. De F-toets zelf blijft overigens in dat hoofdstuk buiten beschouwing.

In het hoofdstuk over regressie en correlatie wordt het erg aardige advies gegeven om het spreidingsdiagram in gestandaardiseerde vorm te tekenen, omdat je dan een indruk kunt krijgen van de samenhang van de variabelen. Bedoeld is natuurlijk de kwaliteit van de (lineaire) samenhang, maar hoe je dat dan doet en waarom het zo is [r = regressiecoëfficiënt maal σ_x gedeeld door σ_y] wordt niet verteld; een gemiste kans.

In het laatste hoofdstuk wordt gesteld dat er een aanmerkelijk verschil is tussen de χ^2 -kansverdeling en de χ^2 -toets. Bedoeld is vermoedelijk dat niet onmiddellijk is in te zien dat de toetsingsgrootheid die ten behoeve van de aanpassingstoets en de toets voor homogeniteit wordt geïntroduceerd, [bij benadering] een χ^2 -verdeling heeft. Maar als dat een didactisch probleem is, is het toch niet zo moeilijk om te laten zien dat de afzonderlijke termen van de berekende toetsingsgrootheid zijn op te vatten als kwadraten van standaard normaal verdeelde grootheden?

Tenslotte een laatste opmerking met betrekking tot de beide toetsen voor gepaarde waarnemingen: de non-parametrische tekentoets en de t-toets. De student moet wel erg goed lezen om zich te realiseren dat bij de overgang

van de eerstgenoemde naar de laatstgenoemde toets niet alleen aandacht wordt geschonken aan de grootte van de waargenomen verschillen, maar dat tegelijkertijd de veronderstelling wordt ingevoerd dat deze verschillen normaal verdeeld zijn. En het is natuurlijk juist deze laatste veronderstelling die ons ervan behoort te weerhouden om klakkeloos de t-toets toe te passen.

Samenvattend: men mag best hopen dat dit boek ook een derde druk zal beleven en men mag verwachten dat die nog wat beter zal zijn dan de tweede.

G.J. van Driel, E.U.R.

Thomas P. Hettmansperger

Statistical inference based on ranks.

John Wiley & Sons, New York, 1984, xvii+323p, ISBN 0-471-88474-X, f 210.-

Verdelingsvrije statistische methoden gebaseerd op rangnummers behoren sinds het eind van de vijftiger jaren tot de standaardtechnieken van de statistische praktijk. Ook theoretisch vormen ze een aardige kluif, omdat bewijzen van asymptotische normaliteit lastig zijn vanwege de afhankelijkheid van de rangnummers waarop de relevante statistische grootheden zijn gebaseerd. Er zijn dan ook al heel wat boeken over dit onderwerp verschenen. Tot de bekendste behoren Lehmann, *Nonparametrics* (1975) en Hájek & Sidák, *Theory of rank tests* (1967). Het eerste is meer elementair van aard en sterk toepassingsgericht (maar bepaald géén kookboek), het tweede is een al wat ouder standaardwerk over de theorie van rangtoetsen.

Het onderhavige boek van Hettmansperger is meer theoretisch dan praktisch gericht. Weliswaar wordt de theorie met voorbeelden geïllustreerd, maar de nadruk ligt toch op het min of meer systematisch ontwikkelen van de theorie van rangmethoden. Daarbij komen de volgende onderwerpen aan de orde: één-, twee, en k-steekproef lokatie problemen, twee-faktor variantie analyse, rangkorrelatie, lineaire regressie modellen en multivariate lokatieproblemen (summier).

Een aantrekkelijk aspect van het boek is dat naast rangtoetsen de korresponderende Hodges-Lehmann schatters veel aandacht krijgen. Daarbij komt ook de robuustheid van schatters ter sprake; de "influence curve" wordt hiertoe ten tonele gevoerd.

De beschouwde toetsen worden onderling systematisch vergeleken m.b.v. de Pitman efficiency. Bewijzen van asymptotische normaliteit van rangtoets-grootheden worden afgeleid uit de projectiestelling van Hájek; soms wordt naar het boek van Hájek & Sidák verwezen. Rang-som toetsen (zoals de toetsen van Wilcoxon) nemen een centrale plaats in. In verband met de regressie-modellen wordt ook enige aandacht besteed aan de asymptotische lineariteit van dergelijke toetsen.

Het wiskundig niveau van het boek is enigszins wisselend; de auteur gaat technische uiteenzettingen niet uit de weg maar maakt géén gebruik van maat- en integratietheorie. Hier en daar is de presentatie wat slordig, zoals bij de bespreking van de Pitman efficiency in sectie 2.6. Sommige geavanceerde onderwerpen, zoals Bahadur efficiency en naburigheid van kansverdelingen met de daarop gebaseerde konstruktie van asymptotisch optimale rangtoetsen, komen niet ter sprake.

Niettemin bevat het boek een grote hoeveelheid up-to-date informatie, die op goed leesbare wijze wordt aangeboden. Een uitstekend studieboek voor eenieder die van dit belangrijke terrein van statistiek wat meer te weten wil komen. De hoge prijs schijnt tegenwoordig onvermijdelijk.

J. Oosterhoff, V.U.A.

P.A. Verheyen, T.M.A. Bemelmans, A. Kapteyn, J. Kriens, P.H.M. Ruys & C.J.C.Th. van Schijndel (red.)

De Praktijk van de Econometrie

H.E. Stenfert Kroese, Leiden/Antwerpen, 1984, ISBN 90 207 1334.5, f 39,50

"De praktijk van de econometrie" is een feestbundel, uitgegeven om een drievoudig jubileum op te sieren: in 1984 werd de vijftiende Tilburgse econometristendag gehouden, werd vijftien jaar geleden de studierichting econometrie ingesteld, en bestond de vereniging van Tilburgse econometrie-studenten vijf jaar. De bundel bestaat uit 31 bijdragen van de hand van in totaal 39 auteurs, die allen iets met de Tilburgse econometrie-opleiding te maken hebben of hebben gehad. Ongeveer de helft van de auteurs is aan de KHT verbonden, de meeste anderen hebben daar gestudeerd.

Twee woorden in de titel kunnen verwarring zaaien: het ene is "praktijk", het andere is "econometrie". Wat betreft de econometrie: de bundel omvat bijdragen over het hele scala van exacte vakken zoals dat in Tilburg wordt gedoceerd, van kwantitatieve bedrijfseconomie via besliskunde, wiskundige economie en informatiekunde tot econometrie in engere zin. Wat betreft de praktijk: in 3 van de 31 bijdragen komen econometrische schattingsresultaten voor, terwijl in 7 andere op de een of andere manier enige meestal illustratief bedoelde cijfers of grafieken voorkomen. Deze wat magere oogst aan empirie suggereert dat er een andere reden moet zijn om de "praktijk" zo te benadrukken. Het zal wel iets met de tijdgeest te maken hebben.

De afzonderlijke bijdragen aan deze bundel zijn uiteraard allemaal aan de korte kant, en lopen uiteen van een beknopte inleiding in nieuwe ontwikkelingen op een bepaald vakgebied tot een samenvatting van recente, gespecialiseerde onderzoeksresultaten of beschrijvingen van praktijkproblemen die met econometrie (in de ruimste zin) zijn opgelost. De artikelen van de eerste soort wisten mijn interesse het meeste te boeien. Dit betrof in het bijzonder enige bijdragen op het gebied van de wiskundige economie, met name de artikelen van Weddepohl en Elbers over het "overlappende-generatiesmodel", en van Withagen over uitputbare natuurlijke hulpbronnen.

Samenvattend: er is een hoop werk verzet om een aardig boekje te maken. Het is geen schande als het in Uw kast ontbreekt, maar het geeft een interessant overzicht over de onderwerpen die anno heden aan een econometrie-opleiding in ruimste zin, althans in Tilburg, aan de orde zijn.

T.J. Wansbeek, R.U.G.

W.R. Dillon & M. Goldstein

Multivariate Analysis methods and applications

John Wiley & Sons, New York, 1984, 587p, ISBN 0-471-08317-8, f 38,85

Dit boek over multivariate analyse is uitgebracht in de Wiley-reeks: Applied Probability and Statistics. De inhoudsopgave toont de vele onderwerpen die in dit boek behandeld worden (tussen haakjes het bijbehorende aantal pagina's).

1. Selected aspects of multivariate analysis (22)
2. Principal components analysis (30)
3. Factor analysis (54)
4. Multidimensional scaling (50)
5. Cluster analysis (52)
6. Multiple regression (42)
7. Some practical considerations: data analysis problems (51)
8. Cross-classified frequency data (58)
9. Canonical correlation analysis (33)

10. Discriminant analysis: the two-group problem (34)
11. Multiple discriminant analysis and related topics
12. Linear structural relations (LISREL) (60)
13. Latent structure analysis (31).

Voorts bevat het boek een appendix I over lineaire algebra en mathematische statistiek en een bescheiden appendix II met statistische tabellen. De lijst van referenties is redelijk omvangrijk (10 pag.). Er zijn geen vraagstukken.

De doelgroep voor dit boek wordt door de auteurs in hun inleiding duidelijk omschreven:

"Books that attempt to cover a large amount of material and forge a path between theory and practice are often found unsatisfactory on both sides. The audience that we wish to reach, however, consists of users who want a sound understanding of the methods being employed without the burden of having to deal with mathematical concepts for which they are not prepared. Wherever possible we have used verbal explanations in place of more direct mathematical arguments. In addition, real data chosen from a variety of disciplines have been employed throughout the text not only to illustrate ideas and methods, but to point out subtleties in the tools".

Naar mijn mening maken de auteurs met dit boek veel van hun doelstelling waar, maar schieten ze ook op niet onbelangrijke punten tekort. Dit boek geeft een goed up-to-date overzicht van de in de multivariate analyse gebruikelijke methoden en technieken. Steeds wordt uitvoerig uiteengezet wat met de verschillende methoden beoogd wordt. Ook wordt uitgebreid aandacht besteed aan de samenhang van de verschillende methoden. Indien nodig wordt een bespreking van beschikbare statistische software gegeven. De vele goede voorbeelden met echte data vervullen hierbij een belangrijke rol. De in deze voorbeelden optredende praktijkproblemen worden uitvoerig en gedetailleerd beschreven en nauwgezet geanalyseerd. Het blijft beslist niet bij een invuloefening van voorafgegeven formules; de interpretatie van de resultaten van de analyse komt uitvoerig aan bod en tal van details worden in deze voorbeelden toegevoegd. Op deze manier wordt zeker een pad gelegd tussen theorie en praktijk.

Een tekortkoming is dat de lezer uit de tekst vaak niet kan opmaken welke formules hij op gezag van de auteurs als juist moet accepteren en welke hij zou moeten kunnen afleiden op grond van voorgaande resultaten als onderdeel van een logisch betoog. De appendix I zal hem hierbij niet helpen (neem b.v. als contrast de eenvoudige afleiding van de normaalvergelijkingen bij lineaire regressie op p.218 en de voor niet-mathematisch geschoolden niet weggelegde afleiding van de likelihoodvergelijkingen bij factoranalyse op p.81, 82). Ook is soms de betekenis van symbolen niet duidelijk of varieert deze per paragraaf.

Ernstiger is echter nog het feit dat vele beweringen of formuleringen strikt genomen onjuist zijn of, erger nog, door hun onduidelijke verbale formulering zich aan een dergelijke kwalificatie onttrekken. Het vermijden van een strikt mathematische betoogtrant door verbale uiteenzettingen is - gelet op de doelgroep - prijzenswaard, maar dit behoeft toch niet slordig te gebeuren! De ingewijde zal zich hieraan slechts ergeren, de eenvoudige lezer zal trachten de formuleringen echt te begrijpen of voor juist aanmerken met alle gevolgen van dien. Een kleine greep:

- p.74: De lezer wordt gewaarschuwd rekening te houden met imaginaire eigenvectoren van symmetrische matrices.
- p.81: (in relatie tot p.61,62): Onduidelijke exercities bij het tellen van parameters in het factormodel. Er had ergens moeten staan: De identificeerbaarheid van de specifieke varianties (Ψ) wordt (hopelijk) gegarandeerd door de Ledermann-conditie $(p-a)^2 \geq p+q$. In dat geval is ook $\Lambda\Lambda'$ (met Λ de matrix van individuele factorladingen) geïdentificeerd. Bij niet-identificeerbaarheid is

verdere analyse zinloos. Voor identificeerbaarheid van Λ is bij identificeerbare Ψ nodig het opleggen van $\frac{1}{2}q(q-1)$ extra (vaak louter technische condities (die niets te maken hebben met de identificeerbaarheid van Ψ), b.v. door te eisen dat $\Lambda'\Psi^{-1}\Lambda$ diagonaal is, hetgeen handig is bij de bepaling van ML-schattingen.

Een dergelijke boodschap die de identificeerbaarheid van de Ψ benadrukt heb ik niet uit de tekst kunnen halen: deze gooit alle parameters op één hoop.

- p.217: Het stochastisch karakter van X kan geregeerd worden als X en ϵ onafhankelijk zijn. De conditie (6.2.10) is niet toereikend. Het vervolg is - i.t.t. wat de tekst suggereert - te begrijpen als X als deterministisch wordt opgevat (In par. 7.1 worden echter weer andere formuleringen gebezigd).
- p.225: Normaliteit van de afzonderlijke ϵ_i impliceert niet de normaliteit van de gehele vector ϵ .
- p.227: Het linker lid van (6.4.11) moet zijn $\text{var}(\hat{Y}_0 - \hat{Y})$
- p.228: De auteurs noemen een voorspellingsinterval ook een betrouwbaarheidsinterval.
- p.271: T.a.v. $b_k = 0/0$ en $\text{var}(b_k) = \infty$: deze onzin kan volgens de auteurs bewezen worden.
- p.295: De toetsingsgrootheid d wordt gedefinieerd als een functie van niet-waarneembare fouten ϵ_t . Er is dus geen toetsingsgrootheid.

Ondanks de tekortkomingen, die de ene lezer meer zullen storen dan de andere, is de aanschaf van dit boek toch beslist de moeite waard voor een ieder die betrokken is bij de concrete toepassing van multivariate technieken.

B.B. v.d. Genugten, KHT

Richard A. Becker & John M. Chambers

S, an Interactive Environment for Data Analysis and Graphics

Wadsworth, 1984, 550p, ISBN 0-534-03313-X, F 131,85

S is een statistisch pakket door AT&T ontwikkeld dat loopt op een beperkte set computers die met het UNIX operating systeem zijn uitgerust. Het boek is allereerst handleiding voor het pakket, maar het kan ook dienen om een indruk te krijgen van de mogelijkheden van S. De leesbaarheid is redelijk, zeker voor diegenen die ervaring hebben met pakketten als SAS of P-STAT en daardoor steeds kunnen vergelijken.

Hoofdstuk 1 en 2 geven een introductie tot S. Opvallend is dat de S-taal het midden houdt tussen een commandotaal en subroutine-calls. De specificatie is kort b.v. $x[\text{row}(x) < \text{col}(x)] + 0$ betekent "maak de bovendiagonaal-elementen van x gelijk aan nul". Daarnaast wordt snel duidelijk dat graphics en nieuwe hardware mogelijkheden een belangrijke rol is S spelen. Interessant is de IDENTIFY functie waarmee nadere informatie over punten in een plot kan worden opgevraagd als men met een interactief beeldscherm werkt.

In hoofdstuk 3 worden basisaspecten van S behandeld. De output van elke S functie is een dataset, dat kan een getal, vector, matrix of algemene output zijn. Bij regressieanalyse onderkent men substructuren als residu-vector en parametervector en daarnaast een "tekstgedeelte" waarin de overige informatie bevat is, zoals de toetsing van de regressiecoëfficiënten. Deze substructuren kunnen via de operatie \$ afgezonderd worden voor verdere bewerking, zoals in $r + \text{reg}(x,y)\$resid$. S probeert van expressies altijd iets zinvols te maken, zonodig worden operaties herhaald. Is y een vector dan worden via $y + y + 1$ alle elementen van y met 1 verhoogd. De operatie $1:7 \times 1:2$ levert de vector 1,4,3,8,5,12,7 op. Het hanteren van matrices vertoont overeenkomsten met PROC MATRIX in SAS.

In hoofdstuk 4 komen de grafische mogelijkheden aan bod. Veel moderne technieken als QQ plots, Box plots en draftsman's plots zijn geïmplementeerd. Superpositie van plots is mogelijk terwijl men met een interactief grafisch beeldscherm kan vastleggen waar men extra tekstinformatie wil invoegen in de plot. De gebruiker krijgt veel mogelijkheden om zijn grafieken naar eigen inzicht op te bouwen. Maar ook onder de standaard omstandigheden zien de grafieken er netjes uit, men krijgt altijd een mooie schaalindeling in ronde eenheden.

Hoofdstuk 5, getiteld Advanced Use of S, behandelt o.a. de apply functie, waarmee S functies op subarrays van arrays kunnen worden losgelaten. Verder de diary functie (logboek) en het datamanagement.

Hoofdstuk 6 is gewijd aan het gebruik van de macroprocessor. Interessant is de mogelijkheid macros die tijdens executie foutlopen vroegtijdig af te breken via de fatal optie. Verder kan men via de optie prompt tijdens uitvoeren van de macro nog informatie toevoeren.

In hoofdstuk 7 dan de dataanalyse. Variantieanalyse technieken en verdeelingsvrije methoden vindt men er niet, wel wordt aangegeven hoe men deze technieken zou kunnen programmeren middels S functies. Bovendien zijn deze technieken in extenso beschikbaar in de standaardpakketten. Regressieanalyse is goed vertegenwoordigd via de functies reg, rreg, lfit en leaps die respectievelijk een kleinste kwadraten regressie, robuuste regressie, l1-regressie en all-possible-subsets regressie uitvoeren. Loglin voert een loglineaire analyse uit, terwijl twoway voor een tweeweg contingentietabel de aangepaste waarden grafisch weergeeft volgens een methode door Tukey bedacht (zie b.v. Exploratory Data Analysis). Van de multivariate technieken zijn aanwezig principale componenten analyse, canonieke correlatieanalyse en discriminantanalyse. Clusteranalyse met het gebruikelijke dendrogram en multidimensional metric scaling zijn beschikbaar. Wat tijdreeksanalyse betreft bestaat de mogelijkheid een economische reeks in trend, seizoen en rest te splitsen, dit alles via nette grafieken opgesierd. Het eind van hoofdstuk 7 laat zien hoe met S Monte-Carlo methoden kunnen worden uitgevoerd. Hier komt de korte taalspecificatie goed van pas. Ook jackknifen en bootstrappen, met standaardpakketten nauwelijks uit te voeren, zijn redelijk simpel te programmeren. Controle van gevoeligheid van de t-toets voor wilde waarden dient als een voorbeeld hoe S ook voor onderzoek van statistische methodieken kan worden ingezet.

Wil men de resultaten van statistische analyses rapporteren waarbij in de rapportage tekst, grafieken, tabellen en formules voorkomen dan wordt in hoofdstuk 8 getoond hoe S samen met de standaard UNIX formatter nroff/troff dit mogelijk maakt. Op dit punt aangeland vraagt men zich af wat dit moois wel niet moet kosten. Het boek besluit met 3 appendices bevattende een documentatie voor de echte gebruiker, een beschrijving van datasets die als oefenmateriaal worden bijgeleverd (natuurlijk ook Iris en Stack-Loss) en een overzicht van de S-functies.

Samenvattend: het boek geeft duidelijk aan wat het doelgebied van S is, n.l. combinatie van dataanalysetechnieken met graphics. S lijkt leemten te kunnen opvullen die de standaardpakketten op dat gebied vertonen (en wellicht zullen blijven vertonen). Het is nuttig te weten dat S als basis heeft gediend voor het expert systeem REX (Regression Expert) dat door AT&T is gebouwd (zie Compstat Proceedings 1984). Daarom is dit boek interessante lectuur voor hen die op de hoogte willen blijven van nieuwe ontwikkelingen op het gebied van statistische software.

H.J.M. Engels, AKZO.

Norman L. Johnson & Samuel Kotz (editors-in-chief)

Encyclopedia of Statistical Sciences

Volume 2, Classification to Eye Estimate

John Wiley & Sons, New York, 1982, viii+613p, ISBN 0-471-05547-6, £ 55,-
(bij intekening op de hele serie: £ 47,- per deel)

Dit is dan deel 2 van een serie van 8 delen, waarvan deel 1 uitgebreid in KM 17 werd besproken.

Aan dit deel (Classification to Eye Estimate) werkten 109 auteurs mee (buiten de redactie), o.a. (we namen een systematische steekproef uit de lijst medewerkers)

- O. Barndorff-Nielsen, Aarhus Universitet, Aarhus, Denmark, over Exponential Families;
- M.H. DeGroot, Carnegie-Mellon University, Pittsburgh, Pennsylvania, over Decision Theory;
- B. Harris, University of Wisconsin, Madison, Wisconsin, over Entropy;
- E.L. Lehmann, University of California, Berkeley, California, over Point Estimation;
- M.B. Rao, University of Sheffield, Sheffield, England, over Damage Models;
- L. Takacs, Case Western Reserve University, Cleveland, Ohio, over Combinatorics.

De 109 auteurs + de redactie nemen in dit deel zo'n 600 trefwoorden voor hun rekening. Een (eveneens) systematische steekproef (op de pag. 11,61,111 etc.) uit de trefwoorden geeft een indruk van de gevarieerdheid van de onderwerpen.

"Cliff-Ord-test", artikel, 5 kolommen, 19 referenties, door P. Burridge;

"Communications in Statistics", artikel, 4 kolommen, 1 referentie, door D.B. Owen;

"Conditional Probability and Expectation", artikel, 19 kolommen, 26 referenties, door Stamatis Cambanis;

"Contingency Tables", artikel, 20 kolommen, 57 referenties, door S.E. Fienberg;

"Cournot, Antoine Augustin", artikel, 2 kolommen, geen referenties, door Sandy L. Zabell;

"D'Agostino test of normality", redactioneel, 2 kolommen, 5 referenties;

"Density Estimation", artikel, 13 kolommen, 38 referenties, door E.J. Wegman;

"Design of Experiments: Bibliography", artikel/bibliografie, 14 kolommen, 1 tabel, 48 geannoteerde referenties, door Gerald J. Hahn;

"Diversity Indices", artikel, 8 kolommen, 18 referenties, door E.C. Pielon;

"Educational Statistics", artikel, 11 kolommen, 36 referenties, door Maurice M. Tatsuoka;

"English Biometric School", artikel, 2 kolommen, 5 referenties, door E. Senneta;

"Estimation, Point", artikel, 18 kolommen, 36 referenties, door E.L. Lehmann;

"Extreme-value distributions", artikel, 15 kolommen, 30 referenties, door Nancy R. Mann & Nozer D. Singpurwalla.

Tenslotte kunnen we zeggen dat ook dit deel, waaraan auteurs van naam en faam hebben meegewerkt, een zeer verzorgde indruk maakt. De hele uitgave is overigens uitstekend verzorgd, ook wat betreft druk etc. en de conclusies uit de recensie van deel 1 van deze serie blijven dan ook recht overeind.

J.J. Dik, Volvo-Car

Onderstaande werken zijn bij het secretariaat van deze boekbesprekingen-rubriek binnengekomen. Indien een lezer er prijs op stelt één of meer van deze werken te bespreken, gelieve hij/zij dit d.m.v. een briefkaartje kenbaar te maken aan de redacteur van deze rubriek. Het adres staat vermeld op de eerste bladzijde van deze rubriek.

CONGRESVERSLAGEN

Statistical Theory and Data Analysis

(Proceedings of the Pacific Area Statistical Conference);

ed. by K. Matusita.

North-Holland Publ.Comp., Amsterdam, 1985, ix+810p, ISBN 0-444-87665-0.

Data Analysis in Real Life Environment

(Ins & Outs of Solving Problems);

ed. by J.F. Marcotorchino, J.M. Proth & J. Jansen.

North-Holland Publ.Comp., Amsterdam, 1985, x+317p, ISBN 0-444-87692-8.

CWI-PUBLIKATIES

J.G.F. Thiemann

Analytic spaces and dynamic programming: a measure theoretic approach

CWI-tract 14.

Centrum voor Wiskunde & Informatica, Amsterdam, 1985, v+96p,
ISBN 90 6196 285 4, f 14,30.

N.M. van Dijk

Controlled Markov processes: time-discretization

CWI-tract 11.

Centrum voor Wiskunde & Informatica, Amsterdam, 1985, i+166p,
ISBN 90 6196 280 3, f 23,80.

W.C.M. Kallenberg, et al

Testing statistical hypotheses: worked solutions

CWI-syllabus 3

Centrum voor Wiskunde & Informatica, Amsterdam, 1985, iii+310p,
ISBN 90 6196 280 3, f 45,10.