

Boekbesprekingen

Redactie:

Arend D. Oosterhoorn
Van Doorne's Transmissie bv
Postbus 500
5000 AH Tilburg
telefoon 013 - 640319

Daroga Singh & F.S. Chaudhary

Theory and Analysis of Sample Survey Designs

John Wiley & Sons, New York, 1987, xi+380 pag., ISBN 0-470-20266-1, £ 24.25

Zo nu en dan verschijnt er een boek over steekproeftechnieken. De opbouw is meestal traditioneel en onderling lijken ze veel op elkaar. 'Theory and Analysis of Sample Survey Designs' vormt hierop geen uitzondering.

Na een inleidend hoofdstuk, waarin basisbegrippen als populatie, steekproef, steekproefeenheid, steekproefkader, steekproef- en niet-steekproeffouten aan de orde komen volgt een uitgebreid hoofdstuk over de enkelvoudige aselechte steekproef. Trekking met en zonder teruglegging, het bepalen van gemiddelden, totalen en percentages en de bijbehorende formules van de variantie worden uitgebreid behandeld. Hierbij passeren alle mogelijke combinaties de revue.

Hoofdstuk 2 toont gelijk de sterke en de zwakke kanten van het boek. Na een korte en duidelijke inleiding worden er met behulp van weinig stellingen en veel corollaries te veel varianten van het behandelde thema weergegeven. Deze corollaries worden in het algemeen niet bewezen. De theorie wordt toegelicht aan de hand van uitgewerkte voorbeelden. Er zijn veel literatuurverwijzingen opgenomen. Deze zijn echter voornamelijk van oudere datum. Slechts zelden wordt literatuur van na 1970 aangehaald. In een voorbeeld over het betrouwbaarheidsinterval wordt vergeten door n te delen.

Hoofdstuk 3 behandelt de gestratificeerde steekproef. Het hoofdstuk gaat o.a. in op constructie van strata en op de verschillende mogelijkheden om de steekproef over de strata te verdelen. Bij de laatste wordt een onderscheid gemaakt tussen de Neyman-allocatie en de optimale allocatie. Verwarrend is dat in een voorbeeld de Neyman-allocatie wordt aangeduid als de optimale allocatie.

De systematische aselechte steekproef wordt behandeld in hoofdstuk 4. In dit hoofdstuk wordt een aardige methode vermeld, die variantieschattingen mogelijk maakt. Helaas bevat de formule een drukfout. De afleiding is voor niet-ingewijden lastig te volgen omdat zonder verwijzing gebruik wordt gemaakt van de formules van Horvitz-Thompson en van Yates-Grundig. In een later hoofdstuk komen deze formules pas aan bod.

Hoofdstuk 5 gaat over het trekken met ongelijke kansen. Een groot aantal trekkingsmethoden worden behandeld. In dit hoofdstuk wordt ook ingegaan op de schattingsmethoden van Des Ray, Horvitz-Thompson en Murthy.

Hoofdstuk 6 gaat over de quotiëntschatter. Hierin worden ook de gescheiden, gecombineerde en multi-variate quotiëntschatter en de productschatter behandeld. De formule van de zuivere quotiëntschatter wordt ontsierd door een drukfout. Vermeld wordt dat door gebruik te maken van de methode van Midzuno de quotiëntschatter zuiver wordt. Vreemd is dat er geen verband wordt gelegd met hoofdstuk 5 waar de methode van Midzuno, onder de naam Midzuno-Sen, zonder referentie naar hoofdstuk 6 beschreven werd. Evenals in andere hoofdstukken staan er in dit hoofdstuk uitgebreide voorbeelden en rekenpartijen.

De regressieschatter met als bijzonder geval de verschilschatter is het onderwerp van hoofdstuk 7.

In hoofdstuk 8 komt de clustersteekproef met gelijke en ongelijke clusteromvang aan de orde. Uitgebreid wordt ingegaan op de efficiëntie van dit ontwerp, compleet met kostenfuncties.

In vervolg daarop behandelt hoofdstuk 9 de meertrapssteekproef. Er worden veel formules weergegeven. Jammer genoeg meestal zonder bewijs en of uitleg. Hierdoor komt het voor de praktijk belangrijke verschil tussen trekken met en zonder teruglegging niet goed uit de verf.

Op de in de praktijk te weinig toegepaste meerfasensteekproef wordt uitvoerig ingegaan in hoofdstuk 10. Het berekenen van de optimale verdeling van de kosten over de verschillende fasen wordt uitgebreid behandeld. De afleiding van de variantieformules krijgt, in afwijking tot eerdere hoofdstukken, uitgebreide aandacht.

Hoofdstuk 11 gaat in op herhaaldelijk onderzoek. Het gebruik van panels en roterende steekproeven valt hieronder. Het schatten van veranderingen in de tijd, van voortschrijdende gemiddelden en van de bijbehorende varianties komt o.a. aan de orde.

Een onderwerp als sequentiële steekproeven waar hoofdstuk 12 over gaat komt men niet vaak tegen in boeken over steekproeven. Na lezing blijkt de theorie in dit hoofdstuk niet veel meer te bevatten dan de techniek van de inverse steekproefmethode.

Het boek eindigt met hoofdstuk 13 over kwaliteit van steekproefonderzoek. Dit hoofdstuk gaat voornamelijk over de zogenaamde niet-steekproeffout.

Het boek bevat veel slordigheden. Essentiële formules bevatten soms zeer storende drukfouten die zonder kennis van zaken moeilijk te doorzien zijn. Het weglaten van de meeste bewijzen maakt het onderkennen van deze fouten extra moeilijk. Verbanden tussen verschillende onderdelen van de tekst worden vaak niet gegeven. Uit de tekst blijkt een grote ervaring en belezenheid van de auteurs. Voor de deskundige die in de theorie thuis is, bevat het boek daarom nog wel wat lezenswaardigs. Het boek bevat een uitgebreide verzameling opgaven. Degene die slechts een enkel boek op het gebied van de steekproeftheorie wil aanschaffen kan echter beter een ander boek kopen.

S.J.M. de Ree

Centraal Bureau voor de Statistiek

J. Hartog, G. Ridder en J. Theeuwes (eds.)

Panel data and labor market studies

North-Holland, Amsterdam, 1990, x + 336 pag.; ISBN 0-444-88461-0, f 175,-

In december 1988 werd in Amsterdam een internationaal symposium gehouden, getiteld 'Panel data and labor market studies'. Het initiatief daartoe werd genomen door de Organisatie voor Strategisch Arbeidsmarktonderzoek (OSA). De achtergrond daarvan is dat OSA zelf panelstudies houdt op het terrein van de arbeidsmarkt. Dit boek bevat de bijdragen aan dat symposium.

De dynamiek van de arbeidsmarkt is in de tachtiger jaren zeer in de belangstelling komen te staan. Voor een goed begrip van wat bijvoorbeeld werkloosheid is en hoe dat bestreden moet worden, is het noodzakelijk inzicht te krijgen in de veranderingen die op de arbeidsmarkt optreden. Het beschikbaar komen van panelstudies en andere longitudinale gegevens hing nauw samen met deze belangstelling.

De bijdragen aan het symposium zijn tamelijk divers; het was dan ook de bedoeling van de organisatoren onderzoekers die actief zijn op het terrein van de arbeidsmarkt of panelstudies (of beide) bij elkaar te brengen. Het leverde vele boeiende bijdragen op. Vele ervan zijn theoretisch van aard; de empirische bijdragen zijn (helaas) in de minderheid. Wellicht dat met dit boek de brug naar de toepassers kan worden geslagen. In het vervolg zal ik in het kort ingaan op de verschillende bijdragen, met name als deze voor de praktijk enige relevantie hebben.

Het eerste deel gaat over meetproblemen. Bound, Brown, Duncan en Rodgers laten zien dat een uiterst belangrijke variabele als de lengte van een werkloosheidsperiode aanzienlijk vertekend gemeten wordt. Korte perioden bijvoorbeeld worden door respondenten niet vermeld. Zij konden dit nagaan door bedrijven waar deze respondenten werkzaam waren, te vragen naar administratieve gegevens. De vraag is of dit verschijnsel in Nederland (waar werkloosheidsperiodes in het algemeen langer zijn dan in de Verenigde Staten) ook zo sterk optreedt. In ieder geval blijkt zorgvuldige meting van fundamenteel belang te zijn. Lancaster en Imbens beschouwen steekproeven uit dynamische populaties; zij gaan na hoe conclusies kunnen worden getrokken over de daklozen in New York als de kans dat iemand in de steekproef van potentiële daklozen belandt, samenhangt met het risico van dakloos zijn (merkwaardig genoeg in deze context wordt dit choice-based sampling genoemd). Ridder bestudeert paneluitval. Het is vrij waarschijnlijk dat uitvallers (eenheden die niet meer aan een panel willen of kunnen deelnemen) geen willekeurige steekproefelementen zijn. Besproken wordt hoe met deze selectiviteit kan worden omgesprongen. Heckman's bijdrage 'A nonparametric method of moments estimator for the mixture of geometrics model' is zeer theoretisch. Mairesse bespreekt de verschillen in schattingen op grond van dwarsdoorsnede- en tijdreeksgegevens, die zelfs blijken te kunnen voorkomen als ze uit hetzelfde panelonderzoek afkomstig zijn; dit wordt geïllustreerd met productiefuncties.

Het tweede deel gaat over een belangrijke klasse van dynamische studies, de duurmodellen. Het centrale begrip in deze studies is de hazard-functie. Kiefer om te beginnen bespreekt methoden om duurgegevens te analyseren als de duurvariabele (bijvoorbeeld werkloosheidsduur) in klassen is onderverdeeld. Wurzel schat werkloosheidsduurmodellen met behulp van gegevens uit het Duitse Sociaal-economisch Panel. Deze modellen zijn volledig parametrisch hetgeen inhoudt dat de hazard-functie vooraf wordt gespecificeerd. De daaropvolgende bijdrage van Narendranathan en Stewart is semi-parametrisch: het verloop van de hazard-functie wordt hier door de data bepaald. Het effect van exogene variabelen als leeftijd en opleiding is in beide benaderingen overigens vergelijkbaar. Ham en Lalonde bespreken een studie waarin wordt nagegaan of een speciaal opleidingsprogramma werklozen sneller aan een baan helpt. In tegenstelling tot andere studies is hier sprake van een experimentele opzet. Toch is ook dan een zorgvuldige statistische analyse nodig om conclusies te kunnen trekken. Met name speelt het

probleem dat de exacte begindatum van de werkloosheidsperiode onbekend is. Veronderstellingen hierover beïnvloeden conclusies over het effect aanzienlijk.

In het derde deel onderzoeken Jensen en Westergaard-Nielsen deense administratieve gegevens over werkloosheidsduren. Een deel van deze duran bestaat uit tijdelijke werkonderbrekingen (temporary layoffs); zij tonen aan dat het belangrijk is deze speciale werkloosheidsperiodes apart te nemen. Vervolgens analyseert Van Ours nederlandse vacatureduren. Zoals werklozen op zoek zijn naar een baan, zijn werkgevers op de arbeidsmarkt actief om hun vacatures te vervullen. Deze complementaire behandeling is interessant omdat het vele onderzoek naar het zoekgedrag van werklozen aantoonde dat de beschikbaarheid van (geschikte) banen zeer belangrijk was bij het vinden van werk.

Het vierde deel gaat over levenscyclus-modellen. MaCurdy bespreekt hoe de veronderstelling van intertemporele substitutie kan worden getest. Deze veronderstelling houdt in dat individuen in de loop van hun leven hun arbeidstijd aanpassen naarmate het loon dat ze kunnen verdienen varieert. Blundell en Meghir gaan in het bijzonder na hoe paneldata gebruikt kan worden voor het schatten van levenscyclus-modellen. Alessie, Kapteyn, Van Lochem en Wansbeek kiezen een iets algemener kader, namelijk het schatten van een stelsel vraagvergelijkingen.

Deel vijf tenslotte bevat bijdragen van Mortensen, Burdett en Spinnewyn. Mortensen stelt de arbeidsmarkt voor als een spel waarin werkgevers lonen aanbieden en arbeiders deze achtereenvolgens beoordelen. Er blijken evenwichtsverdelingen te bestaan van aangeboden lonen enerzijds en reservatielonen anderzijds waarbij niet alleen werklozen maar ook werkenden aanbiedingen kunnen krijgen. Burdett neemt eveneens de vraagkant van de arbeidsmarkt in beschouwing. In het eenvoudige model dat geanalyseerd wordt, blijkt er een unieke evenwichtsverdeling van lonen te zijn; er kan worden nagegaan welke invloed maatregelen als het veranderen van het minimumloon daarop hebben. Spinnewyn tenslotte analyseert hoe werkloosheidsuitkeringen gefinancierd moeten worden en wat de optimale duur en hoogte van uitkeringen is. Ook deze analyse maakt gebruik van een abstract model van de werking van de arbeidsmarkt.

Dirk Gorter

Centraal Bureau voor de Statistiek

Peter Sprent

Applied nonparametrical statistical methods

Chapman and Hall, Ltd., London, 1989, x + 259 pag., ISBN 0-412-30610-7, £ 14.95

Zoals de auteur in zijn voorwoord al in de eerste regel stelt is dit boek vooral een praktische introductie in de niet-parametrische statistiek. Vele voorbeelden worden zeer duidelijk uitgewerkt waarbij ook ruim voldoende aandacht wordt geschonken aan de gemaakte vooronderstellingen. Het boek is niet een volledig handboek voor niet-parametrische methoden maar geeft

een ruim overzicht van de gangbare en een aantal nieuwere methoden. Daar waar methoden niet worden behandeld wordt de lezer verwezen naar recente publicaties.

Hoofdstuk 1 bespreekt in het kort een aantal basis begrippen als schatters, betrouwbaarheidsintervallen, testen van hypothesen en de voor- en nadelen van niet parametrische methoden ten opzichte van parametrische technieken. Verder wordt de permutatie test duidelijk besproken daar deze techniek de basis is van vele methoden. Een verdere bespreking van een aantal statistische begrippen en technieken is gegeven in de appendices 1 t/m 6. Appendix 7 bevat twee sets gegevens welke regelmatig als voorbeeld gebruikt worden en appendix 8 bevat een aantal tabellen. Helaas is de eerste formule welke gebruikt zou zijn om de Tabel A1 te berekenen foutief weergegeven. Of dit te wijten aan twee drukfouten, helaas staan er ook een, gelukkig klein, aantal storende fouten in de voorbeelden, of een slordige notatie is niet na te gaan, maar in appendix A1 wordt de correcte formule gegeven.

Hoofdstuk 2 is volgens de titel gewijd aan lokatie schatters voor een populatie, in feite wordt een aantal tekentoetsen voor de mediaan of mediaan en gemiddelde behandeld en wordt via de hieruit afgeleide betrouwbaarheidsintervallen aandacht gegeven aan de schatters. Voor zowel de toetsen als de betrouwbaarheidsintervallen worden ook de normale benaderingen voor grote steekproeven uitvoerig behandeld. Ook wordt in dit hoofdstuk de tekentoets voor een monotoon verloop geïntroduceerd. Wellicht zou dit beter passen in hoofdstuk 7 dat over bi- en multivariate data handelt. In hoofdstuk 3 komen verdelingstoetsen en rangtransformatie-technieken aan bod. Ook hier weer dezelfde duidelijke aanpak van problemen, welke vooronderstellingen zijn gebruikt en hoe moet het resultaat van de procedure worden gezien. In hoofdstuk 4 worden de toetsen voor gepaarde waarnemingen afgeleid van de reeds eerder behandelde toetsen voor schatters van mediaan of mediaan en gemiddelde van een populatie. De hier als een minder voor de hand liggende toepassing van de tekentoets behandelde McNemar test geeft een goed uitgewerkt voorbeeld van de mogelijkheden van niet-parametrische toetsen.

Hoofdstukken 5 en 6 zullen voor de meeste lezers van groot belang zijn. Deze behandelen namelijk de toetsen voor het vergelijken van populaties, wat toch nog steeds de meest gebruikelijke vraag van onderzoekers is. Hoofdstuk 5 richt zich op het vergelijken van twee populaties, zowel het vergelijken van locatie als van spreiding. Hierbij wordt de Wilcoxon-Mann-Whitney toets zeer uitgebreid en duidelijk behandeld, inclusief hieruit afgeleide schatters en bijbehorende betrouwbaarheidsintervallen. Verder worden de Siegel-Tukey toets en de "squared rank" toets voor verschillen in spreiding en de Smirnov en de Cramer-von Mises toets voor gelijke (cumulatieve) verdelingen behandeld. In hoofdstuk 6 worden voor een aantal behandelde toetsen de uitbreiding naar drie of meer populaties gegeven. De veel gebruikte Kruskal-Wallis als uitbreiding van de Wilcoxon-Mann-Whitney toets, de Friedman toets als uitbreiding van de tekentoets. Voor beide toetsen worden ook de multiple comparison voortzettingen behandeld. Ook worden de analoge toetsen met van der Waerden of "expected normal" scores behandeld.

Voor het toetsen van verschil in spreiding wordt de uitbreiding van de "squared ranks" toets behandeld en voor het vergelijken van (cumulatieve) verdelingen de Kolmogorov-Smirnov toets, beiden echter uiterst summier. Tenslotte wordt de Cochran toets voor binaire data als analogon van de Friedman toets besproken. De volgens de titel aangekondigde behandeling van bi- en

multivariate niet-parametrische technieken in hoofdstuk 7 is vrijwel uitsluitend gewijd aan correlatie (Spearman en Kendall) voor twee variabelen en niet-parametrische regressie met een verklarende variabele. Hoewel deze onderwerpen zeer duidelijk worden behandeld, de auteur blijft naar mijn smaak zelfs te lang stilstaan bij de parametrische kleinste kwadraten methoden, is de behandeling van multivariate technieken teleurstellend. Dit steeds meer interesse krijgend gebied wordt in 1 bladzijde afgedaan. Wellicht is dit onderwerp naar de mening van de auteur te ingewikkeld om in een boek als dit te behandelen. Ook in hoofdstuk 8 wordt een onderwerp behandeld dat volgens de auteur te diep gaat voor dit boek. Toch is de inleiding naar log-lineaire modellen uitstekend zonder al te ver op de theorie in te gaan. Het gehele hoofdstuk is overigens een uitstekende behandeling van het toetsen van categorische data waarbij veel aandacht wordt besteed aan de vooronderstellingen die, veelal onbewust, bij de gekozen toetsen worden gemaakt.

Hoofdstuk 9 behandelt robuuste schatters, jackknife en bootstrap methoden, opnieuw zonder al te diep op de theorie in te gaan. Wel wordt aan de hand van een aantal rekenvoorbeelden de werking van de methoden duidelijk gemaakt. Tenslotte wordt in hoofdstuk 10 iets over meer complexe toetsings- en schattingsmethoden gezegd. Dit is terecht kort gehouden, stipt slechts een aantal methoden aan en geeft voor de technieken uitgebreide literatuur verwijzingen.

Al met al beantwoordt het boek goed aan het in de titel aangegeven doel en is voor zowel statistici als niet-statistici die in de niet-parametrische methoden zijn geïnteresseerd een goede keuze.

J.F.A. Quadt

Unilever Research Laboratorium

A. Gifi

Nonlinear Multivariate Analysis

John Wiley & Sons, Ltd., New York. 1990, xx + 579 pag., ISBN 0-471-92620-5, £ 34.95

In 1980 en 1981 werd door de Vakgroep Datatheorie een cursus Niet-lineaire Multivariate Analyse aangeboden in het kader van het post-doctoraal onderwijs. Centraal in de cursus stonden een intern uitgegeven klapper en een aantal door de medewerkers van de vakgroep geschreven computerprogramma's. De klapper, die in 1980 in het Nederlands verscheen, en in 1981 in het Engels werd herschreven, verscheen onder de auteursnaam Albert Gifi. Albert Gifi was de naam die de medewerkers aan de Vakgroep Datatheorie gebruikten voor hun gemeenschappelijke producten. De schrijvers van de teksten en de computerprogramma's waren toen, volgens de Preface, Bert Bettonville, Eeke van der Burg, John van de Geer, Willem Heiser, Jan de Leeuw, Jacqueline Meulman, Jan van Rijckevorsel en Ineke Stoop. Van deze groep personen zijn nu nog Willem Heiser en Jacqueline Meulman bij de Vakgroep Datatheorie werkzaam, en

zij hebben, met Gerda van den Berg, het boek gereviseerd en voor publicatie bij Wiley gereedgemaakt.

Het voorwoord van het boek is geschreven door Jan de Leeuw (nu werkzaam aan de University of California Los Angeles). Hij was de geestelijke vader van het project, en veel van de ideeën die in Gifi staan zijn reeds te vinden in zijn proefschrift (1973; *Canonical analysis of categorical data*; in 1984 heruitgegeven door de D.S.W.O.-Press). In het voorwoord geeft hij aan waarin het boek uitzonderlijk is. In de eerste plaats gaat het hier om het onderwerp. Het boek gaat over multivariate analyse, maar in het boek staan niet technieken of modellen centraal, maar computerprogramma's. De computerprogramma's berekenen met behulp van het Alternating Least Squares algoritme uitbreidingen van enkele centrale multivariate analysetechnieken. De uitbreidingen passen de multivariate analysetechnieken aan voor de analyse van categorische data. Hierbij wordt een onderscheid gemaakt tussen numerieke categorische variabelen, ordinale variabelen en nominale variabelen. De categorieën in de data matrix worden vervangen door een kwantificatie die mogelijk gerestricteerd is door de eis die het meetniveau (numeriek, ordinaal of nominaal) aan deze kwantificatie stelt. Deze kwantificatie is optimaal in de zin dat een of ander criterium wordt geoptimaliseerd. Zo is het bijvoorbeeld mogelijk om canonische correlatie-analyse te doen met een mengeling van numerieke, ordinale en nominale variabelen. Hierbij worden de kwantificaties zo gekozen dat, bijvoorbeeld, de eerste canonische correlatie is gemaximaliseerd. De technieken (computerprogramma's) die in het boek de meeste aandacht krijgen zijn homogeniteitsanalyse (HOMALS), niet-lineaire principale componenten analyse (PRINCALS), niet-lineaire canonische correlatie analyse (CANALS), niet-lineaire gegeneraliseerde canonische analyse (k-sets analyse; OVERALS) en correspondentieanalyse (ANACOR). Hiervan zijn ANACOR, HOMALS, PRINCALS en OVERALS sinds kort leverbaar als onderdeel van de SPSS module CATEGORIES, wat aangeeft dat de kwaliteit van het werk van Gifi internationaal erkend is. Nog steeds loopt het 'Gifi-project' voort. Het centrale thema lijkt hierbij nog steeds te zijn: bij welke data-analysemethoden kan het idee vruchtbaar worden toegepast om getallen toe te kennen aan categorieën? Deze toekenning vindt steeds plaats met behulp van het Alternating Least Squares algoritme. Recentere producten zijn, bijvoorbeeld, van de hand van Bijleveld (1989; *Exploratory linear dynamic systems analysis*, Leiden: D.S.W.O.-Press), die een computerprogramma voor optimale kwantificatie van categorische lineaire dynamische systemen heeft geschreven (DYNAMALS), en van van Buuren (1990; *Optimal scaling of time series*; Leiden: D.S.W.O.-Press), die vergelijkbaar werk heeft geleverd voor categorische tijdreeksen.

In de tweede plaats is het boek bijzonder in de opvatting die het geeft over multivariate analyse. Gifi besteedt hier veel aandacht aan: de eerste 60 bladzijden van de 500 worden besteed aan de plaatsbepaling van het boek. Dit stuk is bijzonder onderhoudend. Gifi zet zich af tegen de 'mainstream' van toegepaste statistiek: "Of course we have to realize that there is, according to many statisticians, one appropriate way to analyse data: first formulate a model on the basis of prior knowledge, then compute the likelihood of the data given the model, and then estimate and test the appropriateness of the model by maximizing the likelihood function. There are some variations on this scheme, depending on what exactly this prior knowledge includes,

but these variations are minor compared to the prescription seemingly advocated by Gifi. First choose a technique (implemented in a computer program) on the basis of the format of your data, then apply this technique, and study the output. The prior knowledge is not mentioned explicitly, but more importantly it is not isolated in any one particular time and place in the data analysis process. The classical recipe suggests that prior knowledge (theory, prejudice, experience, expectations) enters in the formulation of the model, but then it should not be used at all during the other stages of the process. In the book by Gifi many other places where prior information, and expertise, play a role in data analysis are indicated. Not all thought and creativity of the investigator, and the history of his science, is concentrated in the mysterious phase in which the model is formulated." (in Foreword, p. v,vi, geschreven door Jan de Leeuw; zie ook een eerdere discussie in *Statistica Neerlandica*, 42(2), tussen Molenaar en de Leeuw). Het is zeer onderhoudend te lezen hoe het mijns inziens extreme standpunt van de klassieke toegepaste statistiek met een even extreem standpunt tegemoet wordt getreden!

Ik vind de betogen niet altijd even overtuigend. Om een voorbeeld te geven, het aanvallen van de dominante positie van de multivariaat normale verdeling lijkt me terecht en zeker ter zake, maar aan de andere kant kan hiermee in een boek over categorische multivariate data analyse niet volstaan worden. Hier verwacht ik ergens een opmerking dat de aanname van bijvoorbeeld een multinomiale verdeling in veel toepassingen niet een al te sterke aanname is, maar vind een dergelijke opmerking nergens. In plaats daarvan wordt de statistische benadering van discrete multivariate analyse met loglineaire analyse o.a. aangevallen op het feit dat toetsstatistieken niet meer asymptotisch chikwadraat verdeeld zijn als de frekwenties per cel te laag zijn. Dit mag zo zijn, maar resultaten van Fienberg (1980, appendix) laten zien dat voor conditionele chikwadraattoetsen dit probleem in veel mindere mate geldt.

Enkele andere bijzondere aspecten van het boek zijn de nadruk op kwantificatie van de observaties (rijen) van de data matrix, de nadruk op geometrie in multivariate analyse (hier is de invloed van John van de Geer zichtbaar), de nadruk op stabiliteitstudies middels de jackknife, de bootstrap, kruisvalidatie en randomisatietoetsen (deze methoden werden door Gifi rond 1980 niet alleen gepropageerd, maar ook reeds uitgevoerd) en de aandacht voor voorbeelden.

De hoofdstukken zijn, achtereenvolgens, 2. Coding of categorical data, 3. Homogeneity analysis, 4. Nonlinear principal component analysis 5. Nonlinear generalized canonical analysis, 6. Nonlinear canonical correlation analysis, 7. Asymmetric treatment of sets: some special cases, some future programs, 8. multidimensional scaling and correspondence analysis, 9. Models as gauges for the analysis of binary data, 10. Reflections on restrictions, 11. Nonlinear multivariate analysis: principles and possibilities, 12. The study of stability en 13. The proof of the pudding (een hoofdstuk met voorbeelden). Er is een appendix over gebruikte notatie, een over matrix algebra, en een over kegels en projecties op kegels. Het boek bevat een auteursindex en een onderwerp-index.

Het boek levert bij elk van de behandelde onderwerpen een erg goed en oorspronkelijk literatuuroverzicht (tot 1981). De technieken en computerprogrammas worden evenwichtig behandeld, zonder zich te verliezen in details, maar ook zonder te algemeen te blijven. Er is

ruime aandacht voor voorbeelden. De technische delen van het boek zijn zeker niet gemakkelijk leesbaar, maar op de flap van het boeken lezen we dan ook dat "the book will appeal to graduate students, researchers and practitioners in statistics, psychometrics, biometrics, ecology, regional planning, econometrics, environmental studies and marketing research". Ik zou willen toevoegen: als zij nogal wat tijd willen investeren. Er had wat hier en daar meer zorg aan de presentatie besteed kunnen worden (maar misschien was het boek er dan nooit gekomen): veel van Gifi's presentaties van ideeën, modellen en technieken sluiten niet aan bij bekendere presentaties hiervan. Dit vraagt een extra inspanning van de lezer. Om enkele voorbeelden te geven, bij de behandeling van indicatormatrices wordt niet ingegaan op een mogelijke relatie met het gebruik van dummy-variabelen in multivariate analyse. Bij de bespreking van modellen en technieken (p.33-34) wordt niet gedefinieerd wat nu een model en een techniek is (wordt met model nu bedoeld het model waarmee de data zijn gegenereerd, of het model dat restricties aan de data oplegt; en hoort de keuze voor kleinste kwadraten volgens Gifi nu tot het model, tot de techniek of tot geen van beide?). Hierbij sluit aan dat de gepresenteerde technieken/modellen door middel van een verliesfunctie worden gepresenteerd. Hierbij worden mijns inziens model en optimalisatie van het model onnodig vermengd.

Aan het editen van de editie van 1981 is veel tijd besteed. De tekst is bijzonder mooi gelayout. Inhoudelijk zijn er enkele zaken veranderd, maar het gaat hier volgens de Preface voornamelijk om verduidelijkingen. Er zijn twee uitzonderingen. Het leek de editors beter hoofdstuk 4 van de versie van 1981 als achtste hoofdstuk in de versie van 1990 te plaatsen "for reasons of balance". Gelukkig is deze keuze m.i. niet: nu wordt er in hoofdstuk 6 een correspondentieanalyse uitgevoerd zonder dat correspondentieanalyse behandeld is. Verder hebben de editors er voor gekozen om naar nieuw werk te refereren aan het eind van elk hoofdstuk, in een epiloog. De epiloog verschaft, volgens de preface, "guidance to new developments, especially when it is inspired by - or relevant to - the Gifi system, without trying to be exhaustive". Een dergelijke onderneming is natuurlijk gemakkelijk te bekritisieren. Uitputtend is deze epiloog zeker niet. Ik vond ze vaak onevenwichtig. Ik vroeg me soms af waarom überhaupt naar bepaalde zaken werd gerefereerd, en waarom niet naar andere zaken werd gerefereerd. Bijvoorbeeld, in kruistabelanalyse bestaat momenteel veel aandacht voor maximum likelihood versies van correspondentie-analyse en andere modellen voor optimale schaling (zie bijvoorbeeld Goodman, 1985, AnStat; Gilula and Haberman, 1986, JASA). Hier wordt niet naar gerefereerd. Volgens mij ten onrechte, want juist deze referenties leiden tot grotere bekendheid van correspondentieanalyse onder, bijvoorbeeld, sociologen. Verder is het mij niet duidelijk waarom bij de bespreking van de geschiedenis van homogeniteitsanalyse niet het proefschrift van Van Rijckevorsel (1987) wordt genoemd, die hier een origineel hoofdstuk aan wijdt. En als laatste voorbeeld, het is me niet duidelijk waarom nergens naar Bijleveld's werk aan niet-lineaire dynamische systemen is gerefereerd (de eerste publicatie over DYNAMALS verscheen bij de Vakgroep Datatheorie in 1988). Soms zijn de toelichtingen in de epilogen zo kort dat niet duidelijk is waar de referentie nu precies bijhoort.

Het is jammer dat het zo'n tijd heeft geduurd voordat de gestencilde klapper officieel gepubliceerd werd. Ik denk dat hierdoor het werk van Gifi tot nu toe een veel kleinere invloed

heeft gehad dan zij had kunnen hebben. Om een indruk te geven van de invloed tot nu toe: de gestencilde klappers uit 1980 en 1981 zijn van 1982 tot en met 1989 slechts 4, 8, 10, 8, 12, 12, 15 en 16 maal geciteerd volgens de (Social) Science Citation Index. (Na 8 maanden in 1990 zijn er 7 citaties). In belangrijke artikelen over optimale transformaties in multivariate analyse, zoals het artikel van Breiman and Friedman (1985, JASA; Estimating optimal transformations for multiple regression and correlation), wordt Gifi slechts zelden geciteerd. Dit komt waarschijnlijk ook doordat de Gifi-auteurs niet vaak de moeite namen om hun producten te slijten in belangrijke internationale statistische tijdschriften; uit de referenties van Gifi (1990) blijkt dat, als zij al publiceerden, zij dit meestal in congresboeken of interne rapporten deden.

Het is waarschijnlijk dat het werk van Gifi in de toekomst meer aandacht zal krijgen, nu dit boek in de prestigieuze Wiley-reeks is opgenomen. Het opnemen van enkele computerprogramma's in SPSS zal ongetwijfeld ook leiden tot een grotere vraag naar theoretische verdieping. Gifi (1990) is het werk waarin deze verdieping gevonden kan worden. Het is nog steeds een standaardwerk op haar gebied.

Peter van der Heijden

Vakgroep Empirische-Theoretische Sociologie, Rijksuniversiteit Utrecht.

K.A. Bollen

Structural equations with latent variables

John Wiley & Sons, Ltd., New York, 1989, xiv + 514 pag., ISBN 0-471-01171-1, £ 39.90

In disciplines like the behavioral and social sciences the interest in the analysis of covariance structures is still very strong, despite critical reflections on its methodology (e.g. Freedman, 1987). Ten years ago, no compelling textbook on this subject was available. By now the situation is different, because Bollen has written a rather complete, and thorough textbook on structural equation modeling.

The difficulty level of the book is surmountable. In order to understand most parts of the book, knowledge of the fundamentals of matrix algebra and acquaintance with a basic course in statistics should be sufficient. A review of matrix algebra is included as an appendix though, as is mathematically advanced material such as topics from asymptotic distribution theory.

The material covered by Bollen is too large for a regular one semester course. However, by careful selection it is possible to construct an introductory course, and as easily courses covering more advanced theory.

From a didactical point of view, at an introductory level the textbook of Hayduk (1987) is more attractive than that of Bollen, but the latter covers much more material at an acceptable level for graduate students. Hayduk is extremely limited in the number of models that are exposed: (non-)smoking behavior is an addictive, not a compelling example. Bollen, on the other hand, uses many, different examples from empirical research (social science area,

mainly). Fortunately, for almost any substantial issue an illustration is presented; occasionally, instructive artificial examples are provided.

As to the purpose of the book, the author mentions three aspects: to demonstrate the generality of the LISREL model, to emphasize applications, and to highlight the crucial role played by substantive expertise in most stages of the modeling process.

Chapter 1 gives a clear introduction and a brief historical overview of structural equation modeling, linking regression type models to measurement models, basically. Right from the start Bollen accentuates that the term 'structural' stands for the assumption that the model parameters are not just descriptive measures of association, but rather that they reveal an invariant 'causal' relation. Despite the quotes (disappearing lateron) this is a rather challenging assumption, the pervasiveness of which is discussed at length in Chapter 3.

In Chapter 2 the basic tools for subsequent chapters are launched: model notation, covariances, and the path analysis approach. Attention is being asked for tests and corrections for heteroscedasticity, and for autocorrelated disturbances and errors of measurement in latent variable models; the associated assumptions are seldomly mentioned as explicitly as here. Also, at an early stage the author makes a distinction between observed variables that depend on the latent variables, and those that cause latent variables; clearly a stand is taken in allowing for cause and effect indicators.

The quite interesting problem of causality is treated at length in Chapter 3. Bollen gives the following definition of causality. Consider one variable, say y_1 , which is isolated from all influences except from a second variable called x_1 . If a change in y_1 accompanies a change in x_1 , then x_1 is a cause of y_1 . Next, he claims and discusses that this definition of causality has three components: isolation, association, and direction of influence. It then turns out that perfect isolation is just an ideal, "so it is impossible to make definitive statements about causes" (p. 45). We have to be satisfied with the assumption of pseudo-isolation (the assumption that the disturbance is uncorrelated with the exogeneous variables of an equation). But even so, there are many factors that might threaten such an assumption. "Thus we should recognize the tentativeness of any claims for a causal relation ..." (p. 56). It thus seems that there is much uncertainty left, almost inevitably.

Despite the ambiguities involved, Bollen shows no convinced doubt against the use of causal terminology, and criticisms made by Holland (1986) and Freedman (1987), for example, are only marginally reflected upon. Nor is any reference being made to the discussion of causal issues in social science research, raised by Lieberman (1985). Most attention is spent to the problems of demonstrating the 'age-old issues' of isolation, association, and the direction of causality. The legacy of the causal framework is taken for granted by Bollen. The position taken up is virtually the same as that of Hagenaars (1990); both believe in causal explanations. Hagenaars' main idol is, that it is the most important task of social science to give causal explanations for concrete social phenomena (o.c., p. 9). Thus, the criticism on causal terminology expressed by De Leeuw (1985) also applies here: "It is unfortunate if this terminology is

used freely, but it is even more unfortunate if introductory books do not contain warnings against such practices" (p. 372).

The section on limitations of 'causal' modeling gives a clear presentation of model-data versus model-reality consistency. At the same time, the author is not completely reluctant to consider automatic model-rejection algorithms, like that of Glymour, Scheines, Spirtes & Kelly (1987). In discussing experimental and nonexperimental research, violations of pseudoisolations play a central role.

Despite the slightly untouched controversies, it is advantageous that at least some 40 pages are attached to the problem of causality, where others do not deal with it at all.

Chapter 4 contains an extensive treatment of structural equation models with observed variables (i.e. models for variables assumed to be observed without measurement errors, or path analysis models for directly observed variables). Main topics in this, and similarly in subsequent chapters are: model specification, identification, estimation, and model evaluation.

Relatively much is said about the identification problem and identification checks. However, the introduction of the identification problem lacks clarity. The terminology being used could easily lead to unnecessary confusion: "Thus the goal is to solve for the unknown parameters in terms of the known-to-be-identified parameters" (p. 89). It is also ambiguously defined, what is meant by saying that the population parameters $\text{Var}(x_1)$ or $\text{Cov}(x_1, y_1)$ are identified. As to identification checks, it should be noticed that some recent developments and a review are presented by Merckens (1991), who also built computer programs to automate identification controls.

The ML, ULS, and GLS estimation methods and their main properties are widely covered. In discussing scale invariance and scale freeness the effects of imposing equality constraints to model parameters are neglected. The treatment of topics around intercept estimation and ANOVA models are outdated meanwhile, since in LISREL (Jöreskog & Sörbom, 1988a) as well as in EQS (Bentler, 1989) models with mean parameters (intercept terms, and mean values of latent variables) can be specified easily, without having to perform tricky parametrizations.

Chapter 4 concludes with a number of appendices. First of all, on derivations of the ML fitting functions, and on numerical solutions to minimize fitting functions. Regarding the latter, the book offers the students an illustration rather than insight. Also, in the section of stopcriteria, it is unclear why not to mention the convergence criterion of the absolute value of the first order derivatives of the fitting function with respect to the parameters to be estimated (Jöreskog & Sörbom, 1988a, p. 269, expressions 11.3 and 11.4). Finally, in an appendix an exposure is given of an application of the LISREL and the EQS program. A number of such program examples are given throughout the book. For a recent, thorough review of both most complete structural equation programs the reader is referred to Bollen (1991), who states in an overall assessment of LISREL and EQS: "Given the considerable overlap, we cannot recommend one program over the other" (p. 73). Probably not.

Chapter 5 examines the consequences of random measurement error for univariate and bivariate statistics, multiple regression, and for multiequation systems, under simplifying model assumptions. Results with and without measurement errors are compared.

In an example it is shown that reliability estimates provide a means of estimating error variances that can be incorporated in a model. No profound discussion is offered whether it is allowed to do so, or what the (dis)advantages of such a procedure might be. It should have been noticed, for example, that often full information estimation methods are used. On the other hand, it is satisfying that the author pleads for sensitivity analysis when reliabilities are fixed.

Clearly, however, Bollen does not take as extreme, or outspoken, a view as Hayduk (1987), who finds that measurement reliabilities should routinely be fixed rather than free, and who also thinks that the strategy of fixing measurement error variance at specified values gives the applied researcher "some direct control over the meanings of the concepts" (p. 120). Crossvalidation among researchers would not be bad there, for sure.

Chapter 6 embodies issues of measurement relevant for a structural equation approach to measurement models. After introducing the nature of measurement (the measurement process begins with a concept, of which latent variables are representations), sections follow on validity and reliability concepts, both treated from a traditional (test theoretical) point of view, and within the context of structural equation modeling.

An alternative definition of classical validity is given: "The validity of a measure x_i of ξ_j is the magnitude of the direct structural relation between ξ_j and x_i " (p. 197). It is stressed that the question of validity in a structural equation framework is "whether a causal relation exists between the latent variable and an observed variable" (p. 197). Next, Bollen proposes four measures of validity based on such an approach, among which the unique validity variance, i.e. that part of explained variance in x_i that is uniquely attributable to ξ_j , together with degree of collinearity indicators, seems worth considering.

In a similar spirit an alternative to the classical definition of reliability is proposed: the reliability of x_i is the magnitude of the direct relations that all variables (except error components δ) have on x_i .

Then follows the long, and important Chapter 7 on confirmatory factor analysis, as a procedure to estimate measurement models.

It should be noticed that in the context of assigning to a latent variable the scale of an observed variable, Bollen interprets 'the same scale' of two variables x_i and ξ_j in the rather limited sense of "on average a one-unit shift of ξ_j leads to a one-unit shift in x_i " (p. 240). Little reference is made to the literature here, e.g. on the subject of which indicator to choose in order to solve the scaling problem. Furthermore, it remains uncertain why not to standardize $\text{Var}(u)$ instead of fixing a loading λ_{ij} , as an alternative scaling option in order to have an identified model.

This chapter has an interesting section on identification; even the difference between local and global identification is discussed, and lucidly presented. However, it is unclear why the

problem of empirical identification (e.g. Rindskopf, 1984) is not brought up here, on page 250 for example.

In discussing model evaluation the author demonstrates a careful approach to using chi-square estimates of model fit. For obvious reasons he is rather elaborate on incremental fit indices. However, the interpretation of their size, or how to deal appropriately with these indices, remains a point of great concern. "Awareness of the factors affecting the Δ 's and ρ 's and good judgment are the best guides to evaluating their size" (p. 275). Surely, that is not an easy job to do. Do applied statisticians have good assessment abilities? It could be helpful to use a program like EZPATH (Steiger, 1989, 1990) to compute confidence intervals on population equivalents of the goodness-of-fit index (GFI) and the adjusted goodness-of-fit index (AGFI).

The order of presentation in the book could sometimes be improved. For example in this chapter, while discussing three well-known significance tests for nested models (likelihood ratio, Lagrange multiplier, and Wald) obvious links with the modification index and so-called t-values are not made immediately. Also, with regard to model respecification the concept of model equivalence is not mentioned at all, where that could have been done easily (matrix (7.84) shows identical values for modification indices).

The enumeration of six limitations of exploratory test statistics is very useful. Considering such limitations could in general lead to the necessary cautionary behavior in model modification. Regarding the latter, the author gives the obvious, almost trivial, advices: "Researchers should use the empirical tests as 'advisors' not the 'supervisors' of their model modification", and "cross-validation or replication for an independent sample is an important step in building confidence in the new specification" (p. 305). Such advice is like (redundant) advertisement: if you don't do it, the customers won't buy it, and if you do it, they will probably ignore it.

A recent, updated review (with discussion) on the subject of model evaluation and modification in covariance structure analysis was given by Kaplan (1990).

Chapter 8 contains the basics of the general structural model: the latent variable model and the measurement model are combined to what most readers know as the LISREL model. Apart from topics already covered in preceding chapters for more specific models, attention is paid to the power of significance tests (relatively speaking, a rather weak section; as a matter of fact, Figure 8.9 seems to be the transpose of Table 11.9 in Saris & Stronkhorst, 1984), to the comparison of groups, to missing values (reference to important books like Little & Rubin (1987) and Rubin (1987) is absent), and to the decomposition of effects (very elaborate again).

On page 337 it is suggested that one could minimize the confounding effects of the fitting function employed by checking the residuals associated with different fitting functions. This is doubtful logic. It would hardly be of any help if residuals should diverge for different functions, would it? Its only effect would perhaps be a diminishing confidence in model-data consistency.

In Chapter 9 the following extensions of the general model are treated: alternative model representations (purposefully, without expressing a specific preference for any of them), (in)equality constraints (including tricky specifications), interaction and quadratic terms in a

model, instrumental variable estimators (little on the 2SLS method though, and apparently unaware of the small sample, Monte Carlo results of IV, ML, and ULS estimators for factor analysis, as given by Häggglund (1986)), distributional assumptions (univariate and multivariate normality detection, weighted least squares and elliptical estimators), and categorical variables.

In this chapter much attention is paid to robustness studies concerned with violations of distributional assumptions, and to procedures to correct for negative consequences of such violations. Concise reviews of those studies are presented. In fact, throughout the book robustness questions are treated with great interest. Likewise, in many examples there is ample concern for the detection of outliers and influential observations (Barnett & Lewis (1984) would have been an excellent reference).

The final section on categorical data is instructive, though the scope is too limited to pay full tribute to the potentials of programs like LISCOMP (Muthén, 1988) and PRELIS (Jöreskog & Sörbom, 1988b), programs mainly developed to deal with noncontinuous, categorical and ordinal data.

Some miscellaneous comments

- Most sections are followed by good, concise summaries. That is a relief, in particular after reading sections of the book that are too wordy.
- References to places within the book should have been made by page indication, not by the far less efficient chapter indication.
- Bollen prefers to list the references per author in decreasing order of appearance.
- The total number of errors (typesetting errors and a few serious mistakes) was counted as 34.
- The graphical design and layout of this hard covered book is good, but the layout of tables like 7.5 makes them horrible to read. Also the typesetter was apparently unable to put the hat on θ at its proper place.

It is regarded that Bollen's book is the most extensive and attractive book on structural equation modeling currently available, going beyond a strict introductory purpose. The solidly written book contains a lot of information, and till 1989 say, a fairly complete state of the art of covariance structure analysis for continuous variables.

One major, though hardly escapable, drawback of the book is that the field is still expanding so fast that some of the material was outdated as soon as the author ran to the publisher with the final manuscript. The applied statistician would be most satisfied of course, if the latest theoretical and program developments would be covered by a textbook, but that can barely be demanded nowadays.

If we were his publisher, we would certainly encourage Bollen to make an update for the second edition of the book within due time.

A. Boomsma

Department of Statistics & Measurement Theory
University of Groningen

- Barnett, V. & Lewis, T. (1984). *Outliers in statistical data* (2nd ed.). Chichester: Wiley.
- Bentler, P.M. (1989). *EQS: Structural equations program manual* (Version 3.0). Los Angeles, CA: BMDP Statistical Software.
- Bollen, K.A. (1991). Statistical Computing Software Reviews of EQS (Version 3.0) and LISREL (Version 7.16). *The American Statistician*, 45, 68-73.
- De Leeuw, J. (1985). Review of four books on structural equation modeling. *Psychometrika*, 50, 371-375.
- Holland, P.W. (1986). Statistics and causal inferences. *Journal of the American Statistical Association*, 81, 945-960.
- Freedman, D.A. (1987). As others see us: A case study in path analysis (with discussion). *Journal of Educational Statistics*, 2, 101-128.
- Glymour, C., Scheines, R., Spirtes, P., & Kelly, K. (1987). *Discovering causal structure: Artificial intelligence, philosophy of science, and statistical modeling*. Orlando: Academic Press.
- Hagenaars, J.A.P. (1990). *Idolen van een methodoloog. Over sociaal verklaren in de sociale wetenschappen*. Tilburg: Tilburg University Press.
- Häggglund, G. (1986). *Factor analysis by instrumental variables methods* (Acta Universitatis Upsaliensis, No. 3). Doctoral dissertation, University of Uppsala.
- Hayduk, L.A. (1987). *Structural equation modeling with LISREL: Essentials and advances*. Baltimore: Johns Hopkins University Press.
- Jöreskog, K.G., & Sörbom, D. (1988a). *LISREL 7: A guide to the program and applications*. Chicago, IL: SPSS.
- Jöreskog, K.G., & Sörbom, D. (1988b). *PRELIS: A program for multivariate data screening and data summarization. A preprocessor for LISREL* (2nd ed.). Mooresville, IN: Scientific Software.
- Kaplan, D. (1990). Evaluating and modifying covariance structure models: A review and a recommendation (with discussion). *Multivariate Behavioral Research*, 25, 137-155.
- Liebertson, S. (1985). *Making it count: The improvement of social research and theory*. Berkeley, CA: University of California Press.
- Little, R.J.A., & Rubin, D.B. (1987). *Statistical analysis with missing data*. New York: Wiley.
- Merckens, A. (1991). *Computer algebra applications in econometrics*. Unpublished doctoral dissertation, University of Groningen.
- Muthén, B. (1988). *LISCOMP: Analysis of linear structural relations with a comprehensive measurement model* (2nd ed.). Mooresville, IN: Scientific Software.
- Rindskopf, D. (1984). Structural equation models: Empirical identification, Heywood cases, and related problems. *Sociological Methods and Research*, 13, 109-119.
- Rubin, D.B. (1987). *Multiple imputation for nonresponse in surveys*. New York: Wiley.
- Steiger, J.H. (1989). *EZPATH: A supplementary module for SYSTAT and SYGRAPH*. Evanston, IL: SYSTAT.

Steiger, J.H. (1990). Structural model evaluation and modification: An interval estimation approach. *Multivariate Behavioral Research*, 25, 173-180.

J.C. Gittins

Multi-armed Bandit Allocation Indices

John Wiley & Sons Ltd., New York, 1989, xii + 252 pag., ISBN 0-471-92059-2, £ 29.95

Dit boek behandelt het werk dat Gittins vanaf de zeventiger jaren heeft verricht aan het meer-armige bandiet probleem. Gittins is de specialist op dit onderwerp, en in deze monografie beschrijft hij oude en nog niet eerder gepubliceerde resultaten.

Meer-armige bandiet problemen zijn te beschouwen als een deelklasse van de Markov beslissingsproblemen. Bij een één-armig bandiet probleem zijn er in iedere toestand twee acties: niets doen (de toestand blijft dan ongewijzigd) of continueren (een zekere opbrengst wordt verkregen en de volgende toestand wordt door een kansverdeling gegeven). Een N-armig bandiet probleem bestaat uit N onafhankelijke één-armige bandieten, waarbij op ieder beslissingstijdstip slechts één arm gecontinueerd kan worden en aan de andere armen niets wordt gedaan. het probleem luidt; gegeven de toestand van iedere arm, welke arm kan het beste gecontinueerd worden t.o.v. het optimaliteitscriterium van de verdisconteerde opbrengsten.

Toepassingen van dit model zijn onder andere te vinden op het gebied van de stochastische scheduling, zoekproblemen en in de gezondheidswetenschappen (bijvoorbeeld: welk geneesmiddel kan het beste toegediend worden).

Het hoofdresultaat van Gittins is het bestaan van allocatie indices (door Whittle in zijn voorwoord Gittins indices genoemd), die een zogenaamde index policy als optimale policy genereren. In iedere arm kan voor iedere toestand in die arm een index worden bepaald. Voor het N-armige probleem luidt de oplossing: bepaal in iedere arm de index die bij de toestand van die arm behoort, en continueer de arm met de hoogste index. Gittins bespreekt ook een wat algemener model, waarin per toestand meer dan twee acties mogelijk zijn (superprocessen genaamd). In het algemeen geldt bovenstaand resultaat dan niet meer, maar onder zekere voorwaarden weer wel.

Het boek bestaat uit 9 hoofdstukken. In het eerste hoofdstuk wordt een aantal voorbeelden van meer-armige bandiet problemen gegeven. het tweede hoofdstuk geeft een informeel bewijs en in hoofdstuk 3 wordt de centrale theorie ontwikkeld. Hoofdstuk 4 gaat verder in op de eigenschappen van de indices (bijvoorbeeld monotonie). In hoofdstuk 5 wordt de toepassing van de stochastische scheduling behandeld en komt RESPRO computer pakket ter sprake. Dit pakket moet managers helpen bij de selectie van een bepaald project uit een aantal concurrerende projecten. in hoofdstukken 6 en 7 wordt het probleem als opbrengst proces beschouwd. Met name komen verschillende kansverdelingen (o.a. Bernoulli) voor de opbrengsten aan de orde. Daarvoor kunnen de indices eenmalig worden bepaald en getabelleerd. Hoofdstuk 8 is gewijd

aan zoekproblemen en in het laatste hoofdstuk wordt een alternatief bewijs van de index policy gegeven zoals dat door Whittle in zijn boek "Optimization over Time" (Wiley, 1972) is gedaan.

Al met al is het een boek met enorm veel informatie over bandiet problemen. Het zwaartepunt ligt op de theorie en de toepassing. De algoritmische aspecten voor problemen met eindig veel toestanden blijven wat onderbelicht, wat spijtig is. Het had het boek vollediger gemaakt.

Het is geen eenvoudig boek. Als inleidend boek in dit onderwerp is het tamelijk lastig, zeker als men niet vertrouwd is met zaken als stoptijden. Voor de specialist heeft het heel veel te bieden.

L.C.M. Kallenberg
RU Leiden

Shirley Dowdy & Stanley Wearden
Statistics for Research (2nd edition)

John Wiley & Sons Ltd, New York, 1991, xvii + 629 pag., ISBN 0-471-85703-3, \$ 49.95

Deze uitgave is een uitbreiding van de eerste uitgave met een kleine 100 pagina's. Toegevoegd zijn verdelingsvrije toetsen en bovendien is de uitvoer van statistische software in de theorie opgenomen. Het gebruik van computers bij statistische analyses is tegenwoordig zo algemeen, dat toevoeging van software-uitvoer aan een tekstboek een logisch gevolg is. Een boek dat bestemd is voor onderzoekers rechtvaardigt deze beslissing zeker. Het door de schrijvers gekozen pakket is SAS. Zij motiveren dit door aan te geven dat de in- en uitvoer van de SAS programma's veel overeenkomsten vertonen met die van andere grote pakketten. Bovendien is SAS inmiddels een zeer wijd verbreid pakket geworden.

De aanpak van Dowdy en Wearden maakt op mij een erg doordachte en verfrissende indruk. Door hun jarenlange ervaring met studenten uit een grote verscheidenheid aan research disciplines zijn zij in staat te beoordelen wat in die omgeving leeft. Vanuit deze ervaring hebben zij dit boek geschreven. De opzet van het boek is zo dat per hoofdstuk een vaste indeling gehanteerd wordt die er als volgt uitziet:

- een stuk theorie, afgewisseld met veel voorbeelden
- oefenopgaven, waarbij van een gedeelte van de opgaven (oneven nummers) de oplossingen achterin het boek worden vermeld
- herhaling van een belangrijk stuk theorie in de vorm van een korte procedure beschrijving
- computergebruik met betrekking tot het besprokene; de in- en uitvoer van een SAS programma wordt getoond en van commentaar voorzien
- een serie 'review exercises', die uitsluitend met goed of fout kunnen worden beantwoord (juiste oplossingen achterin het boek)
- een literatuurlijst, met selectie op de besproken onderwerpen

Het boek is geschreven vanuit het gezichtspunt van de research medewerker die wil weten op welke wijze een methode moet worden uitgewerkt, hoe de resultaten moeten worden geïnterpreteerd en die ook begrip wil hebben voor datgene wat hij uitvoert. De nadruk ligt heel duidelijk op het toetsen van hypothesen, in vrijwel ieder hoofdstuk is een onderdeel aan deze techniek gewijd.

In het eerste hoofdstuk wordt een goede motivatie gegeven voor het gebruik van statistische technieken bij onderzoek. Daarna komen in de volgende hoofdstukken een aantal kansverdelingen aan de orde zoals de binomiale verdeling, de Poisson verdeling, de chi-kwadraat verdeling, de normale verdeling en de student's t-verdeling. Zoals al opgemerkt is het toetsen van hypothesen een belangrijke techniek in dit boek. Bij elk van bovenstaande verdelingen wordt een aantal toetsen behandeld, gerelateerd aan de betreffende verdeling. Bovendien wordt de analoge verdelingsvrije toets besproken.

Een zeer groot deel van het boek, meer dan 300 pagina's, wordt besteed aan de onderwerpen variantie analyse, regressie analyse, covariantie analyse en proefopzetten. Dit zijn over het algemeen ook belangrijke onderwerpen voor een research omgeving. Bovenstaande onderwerpen worden breed uitgemeten in dit deel van het boek. Een aantal onderdelen die hier aan de orde komen wil ik er uit lichten.

- Verdelingen van gepaarde variabelen (zeer duidelijk wordt het onderscheid aangegeven tussen regressie en correlatie)
- Enkelvoudige variantie analyse (een goede beschrijving van de multiple comparison procedures wordt gegeven met voor- en nadelen, contrasten zijn opgenomen)
- Variantie analyse model (random en fixed effects, transformaties)
- Andere variantie analyse designs (genest design, randomized complete block design, Latijns vierkant, factorieel design, split-plot design)
- Covariantie analyse
- Multiple regressie en correlatie (matrix procedures, polynomen regressie, orthogonale polynomen)

Een ruime hoeveelheid van 19 tabellen completeert het boek.

De volgende voorbeelden geven aan dat de schrijvers zeer precies te werk zijn gegaan. Bij de introductie van het symbool π als de populatie fractie bij een binomiale verdeling wordt in een voetnoot gemeld dat dit niet het irrationele getal 3.14 voorstelt, opdat de studenten niet in verwarring gebracht worden. Bij de normale verdeling wordt vermeld dat de hier gebruikte π wel het irrationele getal 3.14 voorstelt.

Bij het toetsen wordt het symbool α gebruikt om de fout van de eerste soort aan te geven, meestal op 5% gesteld. Bij regressie analyse wordt de α gebruikt als populatie intercept parameter. Ook hier wordt in een voetnoot aangegeven, dat de gepresenteerde α hier niet het toetsingsrisico voorstelt, maar een populatieparameter in de regressievergelijking, dit eveneens ter voorkoming van misverstanden.

Jammer is het dat de auteurs printerplots hebben gebruikt bij de computeruitvoer en niet de veel fraaiere grafische plots. De plots zijn daardoor bovendien niet altijd even duidelijk. Bij regressie analyse bijvoorbeeld is de regressielijn nu geen getrokken lijn, residupunten vallen weg door het oplossend vermogen etc. Verder mis ik bij de variantie analyse de SAS procedure VARCOMP, terwijl deze procedure voor het schatten van variantiecomponenten juist erg belangrijk is.

Samenvattend, ik vind het een boek boordevol nuttige informatie waardoor het niet alleen als studieboek te gebruiken is, maar ook als naslagwerk goede diensten kan bewijzen. Voor een goed begrip van het boek is een redelijke statistische basiskennis onontbeerlijk, met name wat betreft de toetsingstheorie.

Wim Jansen

AKZO Research Laboratories, Arnhem

Anders Hald

A history of probability and statistics and their applications before 1750

John Wiley & Sons, New York, 1990, xiii + 586 blz., ISBN 0-471-50230-8, £ 55.15.

De schrijver van dit boek is een bekend Deens statisticus. Hij had het plan opgevat om na zijn pensionering een boek over de geschiedenis van de statistiek in de negentiende eeuw te schrijven. Bij de voorbereidingen hiervan bemerkt hij dat de klassieke publikatie van I. Todhunter uit 1865 over de geschiedenis van de waarschijnlijkheidsrekening tot en met Laplace hem niet bevredigde. Zo ontstond dit boek dat kan worden beschouwd als een introductie tot het werk dat hij oorspronkelijk van plan was en dat in 1986 is verschenen onder de titel: *The History of Statistics. The Measurement of Uncertainty before 1900*, geschreven door S.M. Stigler en uitgegeven door The Belknap Press of Harvard University Press.

Het doel van de schrijver was om die bijdragen aan de ontwikkeling van waarschijnlijkheidsrekening en statistiek weer te geven, waarvan gebleken is dat ze van duurzaam belang zijn. Het resultaat is een boek van bijna 600 bladzijden, ingedeeld in 25 hoofdstukken, dat bijna geheel gewijd is aan de periode van 1650 tot 1750. Wat de periode betreft is de titel van werk dus wel enigszins misleidend. Aan bod komen onder andere de bijdragen van Cardano en Galilei, de briefwisseling tussen Pascal en Fermat, de theoretische onderbouwing van de waarschijnlijkheidsrekening door Huygens, verder John Graunt en zijn "Bills of Mortality", levensverzekeringswiskunde waaraan Johan de Witt en Halley, de Moivre en Simpson hebben bijgedragen, het gebruik van wiskundige modellen en statistiek in de astronomie tijdens de wetenschappelijke revolutie, Newtons interpolatie formule, het werk van Bernoulli's, Montmort, De Moivre, de normale benadering van de binominale verdeling zoals geformuleerd door De Moivre enz..

Bij het weergeven van wiskundige gedeelten uit primaire bronnen is gebruik gemaakt van moderne notatie. Volgens de schrijver krijgt de moderne lezer een helder idee van een bewijs, wanneer deze opgeschreven is in de notatie waaraan zij of hij gewend is. Voor deze visie is wel iets te zeggen, met name wanneer de lezer een modern wiskundige is, die niet vertrouwd is met de vroeger gebruikelijke notatie. Echter, in een notatie van een bepaald begrip zit een opvatting daarover. Omdat vroegere en tegenwoordige opvattingen in de waarschijnlijkheidsrekening en statistiek niet hetzelfde zijn, gaan door deze keuze bepaalde historische noties verloren. Verder is het boek duidelijk vanuit het oogpunt van een hedendaags statisticus geschreven, dat blijkt o.a. uit de ruime aandacht die de wiskundige behandeling van de diverse onderwerpen ontvangt.

Hald geeft in het algemeen de belangrijkste wiskundige gedeelten van de relevante werken weer. Soms wordt aan het slot enige aandacht besteed aan de secundaire literatuur over het betreffende onderwerp, maar de nadruk ligt toch heel duidelijk op de reconstructies van de wiskundige bewijsvoering. Dit alles heeft tot gevolg dat het boek eigenlijk alleen interessant is voor moderne statistici en stochastici. Het is indrukwekkend dat iemand in een tijd van ongeveer drie jaar een dergelijk boek kan schrijven. Het is verheugend dat er naast Todhunter een modern wiskundig getinte behandeling van de waarschijnlijkheidsberekening en statistiek van vóór 1900 is verschenen.

De geschiedenis van de waarschijnlijkheidsrekening en statistiek is een onderwerp waarover tegenwoordig veel geschreven wordt. Het merendeel van de lectuur hierover is echter vanuit een ander gezichtspunt geschreven dan het boek van Hald. Aan de titel van één van de meest recente werken is de andere invalshoek af te lezen: *The Empire of Change. How Probability changed Science and Everyday Life*, onder redactie van G. Gigerenzer, Z. Swijtink, T. Porter, L. Daston, J. Beatty en L. Krüger, uitgegeven door Cambridge University Press in 1989.

I. Stamhuis

Vrije Universiteit Amsterdam

O.E. Barndorff-Nielsen en D.R. Cox

Asymptotic Techniques for Use in Statistics

Chapman and Hall, London, 1989, x + 252 pag., ISBN 0-412-31400-2, £ 25.00

Dit is een aardig boek over asymptotische ontwikkelingen van kansverdelingen, die in de moderne mathematische statistiek een grote rol spelen. De inhoud wordt wellicht het best omschreven met de volgende inhoudsopgave:

1. Preliminary notions
2. Some basic limiting procedures
3. Asymptotic expansions
4. Edgeworth and allied expansions

5. Miscellaneity on multivariate distributions
6. Multivariate asymptotic expansions
7. Expansions for conditional distributions.

De auteurs behandelen de stof op de bekende Engelse wijze: zorgvuldige presentatie van de ideeën en technieken, illustratie aan de hand van eenvoudige en minder eenvoudige voorbeelden, en vermijding van bewijzen en technische condities. Stellingen zal men in dit boek slechts bij uitzondering aantreffen!

Het boek lijkt vooral bedoeld voor statistici, die op het betreffende terrein slecht thuis zijn en in kort bestek een goede indruk willen krijgen. De vele voorbeelden kunnen ook in het onderwijs van pas komen. Bovendien wordt relatief veel aandacht besteed aan ontwikkelingen van meerdimensionale verdelingen. Daarentegen ontbreken relaties met de schattingstheorie en de toetsingstheorie bijna geheel. Persoonlijk vind ik dat wat jammer; de nadruk valt nu wel erg sterk op nauwkeurigheid van benaderingen en hoe die te verbeteren. De overlap met een boek als R.J. Serfling, *Approximation Theorems of Mathematical Statistics*, is dan ook verwaarloosbaar.

Het mathematisch niveau is niet onnodig hoog. Met een behoorlijke kennis van de mathematische statistiek en wat elementaire analyse zal men de tekst moeiteloos lezen. De presentatie is helder en het aantal drukfouten klein (al houd ik niet van de notatie $1/2n$ als $n/2$ bedoeld wordt). Een overzicht van verdergaande resultaten en opgaven aan het eind van elk hoofdstuk vergroot de bruikbaarheid.

J. Oosterhoff

Vrije Universiteit Amsterdam

Joseph F. Healey

Statistics : a tool for special research

Wadsworth Publishing Company, Belmont, xiii + 4314 pag., ISBN 0-534-12450-X, £ 15.95

In het voorwoord geeft de auteur aan dat het boek bedoeld is voor studenten in de sociale wetenschappen, met uiteenlopende achtergronden wat wiskunde betreft. Dit laatste is de reden waarom de auteur kiest voor een intuïtieve benadering van de behandelde procedures in plaats van deze te staven met formele wiskundige bewijzen.

Elk hoofdstuk (18 in totaal) begint met een overzichtsschema, wat de student moet helpen de correcte statistische grootheid, toets of procedure te kiezen. Na de tekst wordt ieder hoofdstuk afgesloten met een samenvatting, een lijst met behandelde begrippen, opgaven en een voorbeeld van de behandelde procedure met behulp van SPSS PC+. Van de oneven genummerde opgaven wordt achterin de oplossingen gegeven.

Vooral de lijst met begrippen en de grote hoeveelheid opgaven zullen studenten aanspreken. De voorbeelden uit SPSS PC+ vind ik minder geslaagd, omdat ze te summier zijn, zodat er nauwelijks gesproken kan worden van een dergelijke inleiding in het gebruik van SPSS.

De 18 hoofdstukken zijn verdeeld in 4 grote blokken:

Blok I: Beschrijvende Statistiek (hoofdstuk 1 t/m 5)

Blok II: Inductieve Statistiek (hoofdstuk 6 t/m 12)

Blok III: Associatie (hoofdstuk 13 t/m 16)

Blok IV: Multivariatie Technieken (hoofdstuk 17 en 18).

In blok I worden frequentieverdelingen, centrummaten (modus, mediaan, gemiddelde) en spreidingsmaten (gemiddelde absolute afwijking, range, interkwartielafstand, variantie en standaardafwijking) behandeld, waarbij het meetniveau van de variabele als uitgangspunt genomen wordt. Hoewel de begrippen op een heldere wijze uitgelegd worden, ben ik gestoten op een aantal slordigheden, die bij studenten verwarring kunnen oproepen. Cumulatieve frequenties worden alleen behandeld bij klassifikatie van waarden. Het aantal klassen binnen klassifikatie wordt niet op eenduidige wijze bepaald (b.v. gestart wordt met 10 klassen, later blijken er 13 nodig te zijn). Het klassemidden wordt aangeduid met 'm'. Deze notatie komt terug bij formules voor gemiddelde en variantie:

$$\bar{X} = \frac{\sum fm}{N}$$

Nergens vindt men formules voor het gemiddelde en variantie binnen frekwentieverdelingen, waarbij géén klassifikatie van waarden is toegepast, dus een notatie als:

$$\bar{X} = \frac{\sum fX}{N} \quad (2)$$

In de variantie van de steekproef wordt consequent gerekend met "N" in de noemer, hetgeen in de latere hoofdstukken voor standaardfouten van steekproevenverdelingen en onconventionele notatie tot gevolg heeft. In hoofdstuk 5 worden standaardcodes en de normale verdeling geïntroduceerd. De transformatieformule wordt gegeven, zonder iets te vermelden over lineaire transformaties in 't algemeen en hun effect op gemiddelde en variantie/standaardafwijking. De normale verdeling wordt intuïtief geïntroduceerd (theoretische frequentie verdeling). Tevens worden oppervlakten/kansen bepaald onder de kromme, zonder enige inleiding in kansrekenen.

In blok II worden steekproevenverdelingen en de toetsingstheorie op een heldere wijze uitgelegd. De onderlinge relaties tussen de bijbehorende kansverdelingen ontbreken echter. Behandeld worden de volgende toetsen: één gemiddelde, één proportie, twee (on)afhankelijke gemiddelden, twee onafhankelijke proporties, twee nonparametrische toetsen voor onafhankelijke steekproeven, enkelvoudige variantie-analyse, een toets voor associatie en "goodness of fit" (χ^2)

De keuze van deze toetsen lijkt mij vrij willekeurig: géén toets voor twee gelijke afhankelijke proporties en geen nonparametrische toetsen voor afhankelijke steekproeven. Inconsequente notaties komen voor bij:

- 1) opstellen van hypothesen $H_0: \bar{X} = \mu$, waarmee bedoeld wordt dat een bepaalde steekproef van $\bar{X}$ is uit een populatie met gemiddelde μ ; eenzelfde fout vindt men bij één proportie. Een steekproefkarakteristiek hoort m.i. niet thuis in een hypothese. Deze fout wordt niet gemaakt bij twee of meer steekproeven.
- 2) de verschillende standaardfouten. Zo worden géén apart symbolen gereserveerd voor de standaardfouten van de steekproevenverdelingen van één gemiddelde $\sigma_{\bar{X}}$ en $S_{\bar{X}}$. Wel een apart symbool wordt gebruikt in het twee steekproeven geval: $\sigma_{\bar{X} - \bar{X}}$, echter dit wordt gebruikt voor zowel het geval, waarin de varianties bekend verondersteld worden, als waarin zij onbekend zijn; dit had mijns inziens minstens $\delta_{\bar{X} - \bar{X}}$ of $S_{\bar{X} - \bar{X}}$ moeten zijn. Dezelfde onvolkomenheden vindt men ook terug bij proportie-toetsen. Voorts heeft de

berekening van S (zie boven) notaties als gevolg, zoals: $\sqrt{\frac{S_1^2}{N_1 - 1} + \frac{S_2^2}{N_2 - 1}}$, hetgeen ongebruikelijk is.

Alleen de toets wordt besproken, waarin de varianties gelijk worden verondersteld $\sigma_1^2 = \sigma_2^2$. Naar de toets voor gelijkheid van varianties wordt verwezen als zijnde ANOVA.

Van de nonparametrische toetsen wordt alleen de Mann-Whitney en de Runs toets behandeld. Alternatieven voor twee of meer afhankelijke steekproeven ontbreken. Bovendien wordt volstaan met de normale benadering als toetsingsgrootheid.

Blok III gaat uitgebreid in op associatiematen voor alle meetniveaus (Cramers' V , λ , gamma, Spearmans r_s en Pearsons r en enkelvoudige regressie). Bij de rangcorrelatie van Spearman ontbreekt de formule waarin gecorrigeerd wordt voor geknoopte data.

Blok IV behandelt multiple regressie (zonder toetsen), partiële correlatie en partiële gamma. Gezien deze summiere inleiding in het multivariatie geval was het wellicht beter geweest deze onderwerpen in blok 3 te plaatsen.

Hoewel het boek overzichtelijk en helder geschreven is en bovendien veel opgaven bevat, zou ik er niet toe overgaan het als studieboek voor te schrijven. Met name de onvolkomenheden in het inductieve blok en het ontbreken van achtergrond informatie (kansrekenen, transformaties, relaties tussen kansverdelingen) kunnen alleen verwarrend werken.

W.M. Franken

Faculteit Sociale Wetenschappen, K.U. Nijmegen.

S.P. Gudder

Quantum Probability

Academic Press, Inc., San Diego, 1988, xii + 316 pag., ISBN 0-12-305340-4, \$ 49.50

Quantum Mechanics is een verzamelnaam voor allerlei fysische theorieën, waarmee men het gedrag van elementaire deeltjes (elektronen, protonen, neutronen enz) tracht te verklaren. Uit experimenten blijkt dat bij deze deeltjes interferentieverschijnselen kunnen optreden, zoals bij (monochromatisch) licht door een dunne spleet. Dit dualisme, deeltjes en tegelijk golfgedrag, kan binnen de Klassieke Mechanica niet verklaard worden. Een zeer toegankelijke beschrijving van dit soort verschijnselen kan de lezer vinden in [F]. Om, tot een verklaring van deze verschijnselen te komen zijn verschillende axiomatische theorieën ontwikkeld. De meest gebruikte is misschien wel de axiomatiek ontwikkeld door ondermeer Dirac en Von Neumann, zie bijvoorbeeld [N]. In deze theorie worden de toestanden van een quantum mechanisch systeem beschreven door de eenheidsvectoren van een Hilbertruimte H en de observabelen (plaats, moment, energie, enz.) door zelfgeadjungeerde operatoren op H . De axioma's geven een concreet quantum mechanisch systeem geen uitsluitel over de dimensie van H , omdat deze afhangt van het aantal vrijheidsgraden van het systeem. De verbinding tussen de fysische wereld en deze mathematische structuur wordt gegeven door het volgende axioma:

Als het systeem in toestand ψ verkeert, dan is de kans dat de observabele T een waarde aanneemt in $A \subseteq \mathbb{R}$ gelijk aan $\langle P^T(A)\psi, \psi \rangle$, waarbij P^T de spectraalmaat van T is en waarbij $\langle, \rangle$ het inproduct op H voorstelt.

Door dit axioma wordt de quantum mechanica een probabilistische theorie. Met deze laatste bewering is echter enige voorzichtigheid geboden. In de quantum mechanica kunnen uitkomsten interfereren, dit in tegenstelling tot de op de axioma's van Kolmogorov gebaseerde klassieke kanstheorie.

In het onderhavige boek wordt een benadering van de quantum mechanica gepresenteerd, die gebaseerd is op de taal van de Operationele Statistiek. Deze door Foulis en Randall ontwikkelde theorie werd in een tweetal artikelen in 1972 en 1973 gepubliceerd in het "Journal of Mathematical Physics". Operationele Statistiek is te beschouwen als een enigszins triviale generalisatie van de Kolmogorov opbouw van de kansrekening. In plaats van een uitkomstenruimte beschouwt men een "manual" (kaartenbak) die bestaat uit "physical operations". Een physical operation is op te vatten als een experiment, waarbij een "kaart" hoort met daarop de mogelijke uitkomsten van het experiment. In de klassieke kanstheorie komt een bepaald experiment overeen met een meetbare partitie van de uitkomst ruimte. Omgekeerd is niet iedere manual te beschrijven als een collectie meetbare partities van een of andere uitkomstenverzameling (in de klassieke zin). In het bijzonder is dit niet het geval voor de quantum mechanica. Foulis en Randall ontwikkelen dan een statistiek voor deze gegeneraliseerde kansmodellen.

Hoofdstuk 1 en 2 bevatten overzichten van de Kansrekening en de "Von Neumann Theorie". In hoofdstuk 3 geeft de auteur een samenvatting van de Operationele Statistiek. Hiervan uitgaande, wordt in hoofdstuk 4 een axiomatische model van de quantum mechanica gegeven. De in de axiomatiek van Von Neumann voorkomende Hilbertruimte verschijnt nu als een uit de axioma's afgeleid begrip. Hierdoor wordt de rol, die door de Hilbertruimte gespeeld wordt, verduidelijkt. echter de ontwikkelde theorie voegt niets nieuws toe aan de "Von Neumann Theorie". Vanuit fysisch standpunt heeft dit boek dan ook niet veel nieuws te bieden.

De laatste drie hoofdstukken bevatten toepassingen. Aan de orde komen achtereenvolgens "generalised measure theory", "probability manifolds" en "discrete quantum mechanics".

Het is een zeer goed geschreven boek. Van de lezer wordt enige voorkennis gevraagd op het gebied van Kansrekening. Functionaal Analyse, Klassieke Mechanica (Hamilton en Lagrange formalisme) en Quantum Mechanica.

H. van der Weide

Faculteit technische wiskunde en informatica TUD

Referenties:

- [F] Feynman, R. (1965). *The Feynman Lectures on Physics*. Addison-Wesley, Massachusetts.
- [N] Neumann, J. von (1955). *Mathematical Foundations of Quantum Mechanics*. Princeton University Press, Princeton.

Howell Tong

Non-linear Time Series: A Dynamical System Approach

Oxford University Press, Oxford, 1990, xvi + 564 pag, ISBN 0-19-852224-X, £ 50.00

Het is een opmerkelijk feit dat lineaire tijdreeksmodellen gedurende meer dan een halve eeuw de tijdreeksliteratuur domineren. In het bijzonder heeft dit betrekking op de klasse van stationaire lineaire ARMA modellen met normaal verdeelde storingen. Hun sterke kanten zijn alom bekend. Daarentegen zijn die van niet-lineaire tijdreeksmodellen tot voor kort wat onderbelicht gebleven. Met name kunnen deze laatste modellen, in tegenstelling tot de lineaire modellen met normaal verdeelde storingen, belangrijke karakteristieken in tijdreeksen weergeven. Voorbeelden zijn de tijdomkeerbaarheid van een proces, tijdreeksen met niet-normaal verdeelde marginale verdelingsfuncties, en tijdreeksen die op toevallige momenten uitbarstingen van erg grote amplitude vertonen.

Recentelijk valt in de literatuur een toenemende belangstelling te constateren voor niet-lineaire tijdreeksanalyse. Kenners weten dat professor Tong tot een van de voortrekkers op dit gebied

behoort. Het thans verschenen boek toont dit nog eens overduidelijk aan. Het boek is opgebouwd uit de volgende hoofdstukken:

1. Introduction
2. An introduction to dynamical systems
3. Some non-linear time series models
4. Probability structure
5. Statistical aspects
6. Non-linear least-squares prediction on non-linear models
7. Case studies

Voorts zijn drie appendices aan het boek toegevoegd, te weten: 1. Deterministic stability, stochastic stability and ergodicity; 2. Martingale limit theory en 3. Data.

Hoofdstuk 1 bevat een aantal basisbegrippen uit de lineaire tijdreeksanalyse. Tevens worden enige voor- en nadelen van ARMA modellen besproken. In hoofdstuk 2 wordt uitvoerig ingegaan op diverse ontwikkelingen in de lineaire- en niet-lineaire oscillatietheorie. De volgende onderwerpen komen daarbij ondermeer aan de orde: limiet cycli, Volterra expansie, chaos, continue dynamische systemen uitgedrukt in niet-lineaire differentiaalvergelijkingen, niet-lineaire differentievergelijkingen, en stabiliteitstheorie van differentievergelijkingen. Hoofdstuk 3 behandelt een aantal niet-lineaire tijdreeksmodellen. Ook wordt kort ingegaan op de historische achtergrond van enkele modellen. Begrippen zoals deterministische- en stochastische stabiliteit, ergodiciteit, inverteerbaarheid en tijdomkeerbaarheid komen ter sprake in hoofdstuk 4. Hoofdstuk 5 begint met een reeks van nuttige suggesties om met behulp van grafische technieken niet-lineariteiten in data op te sporen. Daarnaast passeren diverse toetsen op lineariteit de revue, en worden enkele schattingstechnieken besproken. Ondanks de omvang van dit hoofdstuk (130 pagina's) wordt de vaak complexe theorie slechts sporadisch toegelicht aan de hand van voorbeelden. In hoofdstuk 6 (10 pagina's) laat de auteur zien hoe men in principe via numerieke integratie voorspellingen met niet-lineaire AR modellen kan genereren. Qua inhoud en diepgang is dit hoofdstuk mijns inziens ronduit teleurstellend. Klaarblijkelijk is dit terrein nog nagenoeg onontgonnen. De niet-triviale mathematisch-statistische problemen verbonden aan het vaststellen van noodzakelijke en voldoende voorwaarden voor de inverteerbaarheid van een niet-lineair model, hetgeen een vereiste is om daarmee te voorspellen, kunnen hieraan debet zijn. Tenslotte worden in hoofdstuk 7 een viertal case-studies gepresenteerd van niet-lineaire tijdreeksen, waaronder de reeds in talloze artikelen geanalyseerde Canadese lynx data en de evenzeer overbekende zonnevlekken van Wolf.

Indien men zich wil verdiepen in de huidige stand van zaken danwel op zoek is naar potentiële onderzoekproblemen op het terrein van de niet-lineaire tijdreeksanalyse, kan ik dit boek warm aanbevelen. Wel dient de lezer te beschikken over een behoorlijk niveau aan kennis op het terrein van de mathematische-statistiek en de stochastiek. Toegepaste statistici zullen minder uit de voeten kunnen met de inhoud van het boek. Helaas schieten de oefeningen die elk hoofdstuk afsluiten naar mijn mening hun doel voorbij daar zij in het algemeen een grondige

kennis van de vaak zeer theoretische literatuur vereisen. Daarnaast zijn de "uitgeputte" case-studies in hoofdstuk 7 weinig behulpzaam om de diverse praktische aspecten van het identificeren, schatten en voorspellen van niet-lineaire tijdreeksen nader toe te lichten, laat staan om deze aspecten enigszins in de vingers te krijgen.

Jan G. de Gooijer
Universiteit van Amsterdam

> > > > > > > > >

Binnengekomen boeken

< < < < < < < < <

Om de nieuws waarde van de boekbesprekingrubriek te verhogen zal in ieder nummer een lijst worden opgenomen met daarin vermeld de bij de redactie binnengekomen boeken. Omdat er soms nogal wat tijd kan zitten tussen de uitgave van het boek en de publicatie van de recensie, hopen we hiermee de dienstverlening aan de lezers te vergroten. Niet alle boeken in deze rubriek zijn al ondergebracht bij een recensent. Mocht er bijzondere belangstelling bestaan voor het bespreken van een boek, dan kan contact op genomen worden met de redacteur. Met nadruk wordt gesteld dat niet zeker is of het boek nog 'in de aanbieding' is.

Anderson, T.W. & K.B. Athreya & D.L. Iglehart

Probability, statistics and mathematics, papers in honor of Samuel Karlin

Academic Press, Inc., San Diego, 1991, xl + 371 pag., \$ 39.95

Mayer-Wolf, E. & E. Merzbach & A. Shwartz (eds)

Stochastic analysis

Academic Press, Inc., San Diego, 1991, xx + 532 pag., \$ 69.95

Moors, J.J.A.

Statistiek in de economie twee, steekproefmethoden en analyserende statistiek

Academic Service, Schoonhoven, 1991, x + 340 pag., f 43.50

Moors, J.J.A.

Statistiek in de economie één

Academic Service, Schoonhoven, 1991, xii + 415 pag., f 43.50

Harter, H.L.

The chronological annotated bibliography of order statistics, Volume III, 1960-1961

American Science Press, Inc., Syracuse, 1991, vi + 214 pag., \$ 95.00

Hendriks, Th.H.B. & P. van Beek

Optimaliseringstechnieken, principes en toepassingen (3e druk)

Bohn Stafleu Van Loghum, Houten, 1991, ix + 348 pag., f 49.00

Snell, E.J. & H.R. Simpson

Applied statistics, a handbook of Genstat analysis

Chapman & Hall, London, 1991, xii + 132 pag., £ 18.50

Stuart, A. & J.K. Ord

Kendall's Advanced theory of statistics volume 2

Edward Arnold, 1991, xii + 719 pag., £ 60.00

Verschuren, P.J.M.

Structurele modellen tussen theorie en praktijk

Het Spectrum, Utrecht, 1991, 662 pag., f 79.90

Fleming, T.R. & D.P. Harrington

Counting processes & survival analysis

John Wiley & Sons, Ltd., New York, 1991, xiii + 429 pag., £ 39.80

Jackson, J.E.

A user's guide to principal components

John Wiley & Sons, Ltd., New York, 1991, xvii + 569 pag., £ 55.15

Levy, P.S. & S. Lemeshow

Sampling of populations, methods and applications

John Wiley & Sons, Ltd., New York, 1991, xxiv + 420 pag., £ 39.80

Morgenthaler, S. & J.W. Tukey

Configural polysampling, a route to practical robustness

John Wiley & Sons, Ltd, New Yk, 1991, xiv + 228 pag., £ 32.15

Pilz, J.

Baysian estimation and experimental design in linear regression models

John Wiley & Sons, Ltd., New York, 1991, x + 296 pag., ?

Balin, L.J. & M. Engelhardt

Statistical analysis of reliability and life-testing models 2nd edition

Marcel Dekker, Inc., New York, 1991, vii + 496 pag., \$ 132.25

Cryer, J.D. & R.B. Miller

Statistics for Business, data analysis and modelling

PWS-Kent Publishing Company, 1994, xx + 811 pag., ?

Hutchinson, T.P.

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

Rumsby Scientific Publishing, Adelaide, 1991, xiv + 266 pag., \$ 24.00

Arnold, B.C. & N. Balakrishnan

Relations, bounds and approximations for order statistics

Springer-Verlag, Berlin, 1991, x + 173 pag., DM 36.00

Benveniste, A. & M. Métivier & P. Piouret

Adaptive algorithms and stochastic approximations

Springer-Verlag, Berlin, 1991, xi + 365 pag., DM 114.0

Brockwell, P.J. & R.A. Davis

ITSM: an interactive time series modelling package for the PC

Springer-Verlag, Berlin, 1991, xi + 104 pag., DM 98.00

Brockwell, P.J. & R.A. Davis

Time series: theory and methods second edition

Springer-Verlag, Berlin, 1991, xvi + 577 pag., DM 98.00

Gerber, H.U.

Life insurance mathematics

Springer-Verlag, Berlin, 1991, xiii + 131 pag., DM 98.00

Lancaster, H.O.

Expectations of life

Springer-Verlag, Berlin, 1991, xv + 605 pag., DM 320.00

Oksendal, B.

Stochastic Differential Equations, an introduction with applications

Springer-Verlag, Berlin, 1991, xv + 186 pag., DM 48.00

Salomom, M.

Deterministic lotsizing models for production planning

Springer-Verlag, Berlin, 1991, vii + 158 pag., dm 46.00

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

Statistical methods: the geometric approach

Springer-Verlag, Berlin, 1991, xv + 560 pag., DM 98.00

Sen, A. & M. Srivastava

Regression analysis, theory, methods, and applications

Springer-Verlag, Berlin, 1991, xv + 347 pag., DM 88.00

>>>>>>>>>>

CWI publicaties

<<<<<<<<<<<

A.M.H. Gerards

Graphs and polyhedra: Binary spaces and cutting planes

CWI, Amsterdam, Tract 73, 1991, iv + 188 pag., ISBN 90-6196-390-7, f 49.00

>>>>>>>>>>

ILO publicaties

<<<<<<<<<<<

Prokopenko, J. and I. Pavlin

Entrepreneurship development in public enterprises

International Labour Organisation, Geneve, 1991, ii + 208 pag., ISBN 92-2-107286-X,
Sfr 27.50

Husmanns, R., F. Mehran and V. Verma

Surveys of economically active population, employment, unemployment and underemployment: An ILO manual on concepts and methods

International Labour Organisation, Geneve, 1991, xvi + 409 pag., ISBN 92-2-106416-X,
Sfr 40.00

