

KM 16(1984)
pag 177-192

BOEKBESPREKINGEN

Redactie: J.J. Dik

Correspondentie adres: Subfaculteit Wiskunde UvA,
Roetersstraat 15
1018 WB Amsterdam
(020) 522-3090

William G. Cochran

Planning & Analysis of Observational Studies

John Wiley & Sons, 1983, 145p, ISBN 0-471-88719-66, £ 20.85

Dit boek handelt over problemen die tijdens de planning- en analysefase van vergelijkende observationele studies kunnen optreden. Veel aandacht wordt besteed aan het effect op de interpretatie van de resultaten, als verkeerde assumpties gesteld worden bij het hanteren van statistische procedures.

Het niveau van de behandelde stof is afgestemd op aankomende wetenschappers uit de sociaal- en gedragswetenschappelijke hoek. Ook beginnende onderzoekers, die aangewezen zijn op observationeel onderzoek, zullen in dit boek een goede inleiding in die materie vinden. Men spreekt van observationeel of quasi-experimenteel onderzoek, wanneer de indeling in groepen naar een bepaald kenmerk of behandeling tot stand komt door:

- zelf-selectie (bijv. dragers versus niet-dragers van gordels)
- natuurlijke selectie (bijv. te vroeg geboren versus op normaal tijdstip geboren babies)
- toedoen van anderen dan de onderzoeker (bijv. verschillende vormen van operaties)

Vaak is men geïnteresseerd in het schatten van een bepaald behandelingseffect, of in de verschillen tussen twee of meer soorten behandelingen. De schrijver laat voor verschillende situaties zien welke storende factoren aanleiding kunnen geven tot misleidende conclusies (bias effect). De meest voorkomende problemen in iedere fase van het onderzoek worden op een heldere wijze in aparte hoofdstukken behandeld.

Hoofdstuk 1: Variation, Control and Bias, handelt over het begrip observationeel onderzoek.

Hoofdstuk 2: Statistical introduction, behandelt onderwerpen die bij de interpretatie van de resultaten een belangrijke rol spelen. Standaard technieken voor toetsen en schatten worden op een elementair niveau besproken. Met behulp van eenvoudige voorbeelden laat de schrijver zien hoe onrechtmatig gebruik van deze technieken kan leiden tot absurde uitspraken. In het tweede gedeelte van dit hoofdstuk worden illustraties gegeven van situaties waarin onbedoelde systematische verschillen tussen populaties kunnen ontstaan. Duidelijk wordt gemaakt dat enige kennis over de richting en grootte van de bias, soms het gebruik van standaard technieken voor praktische doeleinden zinvol maakt.

Hoofdstuk 3. Preliminary aspects of planning, handelt over problemen aan het begin van de planningsfase, zoals de formulering van de vraagstelling, mogelijke bronnen van bias, tijdstip van meting etc.

Hoofdstuk 4: Further aspects of planning, in dit hoofdstuk worden formules gepresenteerd om de steekproefgrootte te schatten die nodig is voor adequate

toetsen voor significantie en schattingsprocedures. Eenvoudige voorbeelden verduidelijken het nut van deze formules. De schrijver schenkt bovendien ruime aandacht aan problemen bij het toepassen van deze formules in observationeel onderzoek.

Hoofdstuk 5: Matching, handelt over technieken om de bias, veroorzaakt door "confounding variables" te reduceren. Matching is een techniek die gebruikt kan worden tijdens het verzamelen van gegevens. Aan de orde komen: within-matching, caliper matching, nearest available matching en mean matching.

Hoofdstuk 6: Adjustments in analysis, handelt over statistische technieken die tijdens de analysefase rekening houdt met confounders. Zowel in een experimentele als in een observationele setting worden matchings- en regressietechnieken met elkaar vergeleken.

Hoofdstuk 7: Simple study structures, bespreekt de structuur van enkele eenvoudige designs die regelmatig in de medische en sociale wetenschappen gebruikt worden.

E. Tan, R.U.Limburg.

M. Hazewinkel & A.H.G. Rinnooy Kan (ed.)

Current Developments in the Interface: Economics, Econometrics, Mathematics

D. Reidel Publishing Company, Dordrecht, 1982, X+355p, ISBN 90-277-1505-X

Dit boek is het verslag van een symposium dat van 12 tot en met 15 januari 1982 te Rotterdam is gehouden ter ere van het 25-jarig bestaan van het Economisch Instituut van de Erasmus Universiteit. Dertien top-mensen uit de economische professie zijn daarbij uitgenodigd een lezing te houden. Het boek bestaat dan ook uit dertien hoofdstukken die elk bestaan uit een introductie van de spreker, het gepresenteerde paper en een discussie. Deze dertien papers zijn:

- 1) H. Theil: "Research in econometrics and the economics of academia". Dit is een anekdotisch verhaal over Prof. Theil's Rotterdamse jaren, gevolgd door een beschouwing over het Amerikaanse universitaire systeem.
- 2) J.H. Drèze: "Decision criteria for business firms". Gegeven het feit dat de winstgevendheid van alternatieve productieplannen van ondernemingen afhangt van een aantal onzekere factoren, wordt de vraag gesteld welke keuze uit de verzameling van mogelijke productieplannen optimaal is. Daarbij wordt speciaal aandacht besteed aan de rol van aandelen koersen als informatiebron voor de ondernemer bij het nemen van deze beslissing.
- 3) J. Durbin: "More than twenty-five years of testing for serial correlation". Dit artikel beschrijft hoofdzakelijk de eigen bijdragen van Prof. Durbin op dit gebied, en wel met name zijn recente publicaties en ideeën over het toetsen van hogere orde autocorrelatie in regressiemodellen met niet-stochastische regressoren. Deze laatste onrealistische, veronderstelling doet afbreuk aan de relevantie van het betoog.
- 4) D. Gale: "Efficiency". Het artikel behandelt op heldere wijze de vraag onder welke voorwaarden allocaties van goederen, winsten e.d., over economische agenten efficient (in de betekenis van Pareto-optimaal) zijn.
- 5) L. Johansen: "Economic models and economic planning and policy: some trends and problems". Dit is een zeer leesbaar artikel over de constructie en het gebruik van macro-econometrische modellen en de wijze waarop nieuwe economische inzichten zoals "supply side economics" en "rational expectations" kunnen worden geïncorporeerd.
- 6) D.W. Jorgenson: "An econometric approach to general equilibrium analysis". Het artikel beschrijft een empirisch multi-sectoraal productiemodel tezamen met een stelsel van vraagvergelijkingen. Een en ander vormt de aanzet tot een empirisch algemeen evenwichtsmodel.
- 7) R.E. Kalman: "Identification from real data". Dit is een gepassioneerd betoog waarin de auteur zich afvraagt hoe lineaire relaties tussen variabelen kunnen

worden geïdentificeerd op basis van hun covariantie matrix.

8) L. Kantorovich: "Planning, mathematics and economics". Dit artikel handelt over de rol van mathematische programmeringsmethoden en schaduw prijzen in de planning van de Sovjet economie.

9) D. Kendrick: "Stochastic control and uncertainty in dynamic economic systems". Het artikel handelt voornamelijk over mogelijk vervolgonderzoek op het gebied van stochastische "control"-theorie. Het is daarom weinig concreet.

10) E. Malinvaud: "An economic model for macro-disequilibrium analysis". De auteur illustreert aan de hand van een eenvoudig macro-economisch onevenwichtigheds-model welke schattingsproblemen zich kunnen voordoen en op welke wijze die problemen op te lossen zijn.

11) R.K. Mehra: "Identification in control and econometrics". Het artikel handelt over het schatten van multivariate tijdreeksen modellen op basis van een zg. "State-space" methodologie. Voorts worden een aantal voorspel experimenten besproken.

12) M.J.D. Powell: "Algorithms for constrained and unconstrained optimization calculations". Dit is een handzaam overzichtsartikel over de huidige stand van zaken met betrekking tot numerieke optimaliseringsprocedures.

13) C.A. Sims: "Scientific standards in econometric modelling". Dit artikel geeft een herbezinning op wat econometristen doen en zouden moeten doen. De auteur benadrukt daarbij het onderscheid tussen subjectiviteit en objectiviteit bij het modelleren van economische data. Het betoog is nogal compact, maar wel de moeite van het lezen waard.

Helaas hebben slechts enkele van de bovengenoemde artikelen (Gale, Johansen, Kantorovich, Powell) min of meer het karakter van een overzichtsartikel. De overige artikelen, met uitzondering van Theil's bijdrage, zijn specialistisch van aard. De bijdrage van Prof. Theil vormt een dissonant in dit boek. Zijn verhaal is meer geschikt voor de weekendbijlage van een dagblad dan voor een bundel wetenschappelijke verhandelingen.

H.J. Bierens, UvA.

C. Bongers

Standardization - Mathematical Methods in Assortment Determination

Martinus Nijhoff Publishing, Boston, 1980, VIII+248p, ISBN 0-89838-029-4, \$ 25,-

Dit boek is de handelseditie van het proefschrift waarop de auteur in 1979 aan de Erasmus Universiteit promoveerde bij Prof. J. Sittig.

Het onderwerp is normaliseren van maten. Voor de producent liggen de voordelen onder andere in lagere produktiekosten: langere series, eventueel minder produktielijnen, minder kosten van distributie. Daar staat tegenover dat de consument niet altijd kan krijgen wat hij precies wil hebben, waardoor er economische en/of subjectieve aanpassingsschaden kunnen ontstaan. Veel gebruikte normalisering- en bestaan uit rekenkundige of meetkundige reeksen; bij de samenstelling ervan worden economische overwegingen niet expliciet gebruikt. De auteur van dit boek doet dit wel en tracht de verwachte aanpassingsschaden van de consumenten te minimaliseren. Deze schaden hangen af van de vraagverdeling naar de verschillende maten, de schaden bij een bepaalde aanpassing en het te kiezen matenpatroon. Vraagverdelingen en schaden moeten uit waarnemingen geschat worden, het te kiezen matenpatroon biedt gelegenheid tot optimaliseren.

De bestudeerde problemen beperken zich tot situaties waarin er het zij één, het zij twee grootheden genormaliseerd moeten worden, de maten niet additief zijn, noch gemakkelijk gesplitst kunnen worden tot de gewenste maten en tenslotte vastgesteld kunnen worden onafhankelijk van andere bestaande matenpatronen. Als x de gewenste maat is en t de gekochte, dan worden schadefuncties van de vorm

$$v_1(t, x) = k_1 (x^\alpha - t^\alpha)^\beta \quad (x \geq t \geq 0),$$

respectievelijk

$$v_2(t, x) = k_2 (t^\alpha - x^\alpha)^\beta \quad (t \geq x \geq 0),$$

$\alpha, \beta \geq 0$, gehanteerd. De beschouwde vraagverdelingen zijn de homogene, de normale en de lognormale verdeling, respectievelijk de bivariate vormen daarvan. Uitgaande van deze gegevens is het mogelijk optimale normalisatiepatronen te berekenen via zogenaamde recursieformules. Deze worden eerst algemeen afgeleid, waarbij wordt aangegeven hoe de berekening met behulp van het zij intervalsectie, het zij de methode van Newton kan geschieden. Daarna volgen formules voor de genoemde kansverdelingen. In het geval van twee te normaliseren afmetingen is een verdere beperking van het keuzepatroon noodzakelijk om redelijk hanteerbare modellen te verkrijgen. Hiervoor worden achtereenvolgens de roosterrestrictie, de lijnrestrictie en de gegeneraliseerde roosterrestrictie gebruikt. Zowel het geval van één afmeting als dat van twee afmetingen wordt geïllustreerd met een gedetailleerd uitgewerkt voorbeeld. Het laatste voorbeeld betreft dikte en breedte van plaatstaal; de resultaten worden vergeleken met de bestaande Euro-norm.

De stof is verdeeld over 18 korte hoofdstukken en 4 appendices.

De studie betreft een interessant maatschappelijk probleem. De uitwerking wekt de indruk dat de zaak van alle kanten bekeken is. Een zekere wijdlopigheid is hiervan het gevolg. Het boek is helder geschreven. De vereiste voorkennis beperkt zich tot enkele grondbegrippen van de statistiek en standaardmethoden om extremen van meerdimensionale functies te bepalen. Doordat het gehele normaliseringsproces bekeken wordt, passeren ook enkele bekende methoden op statistisch en numeriek wiskundig gebied, de revue, toegespitst op de behandelde problemen. Daarbij steekt de zorgvuldige verwijzing naar statistiek literatuur merkwaardig af tegen het geheel ontbreken van literatuurverwijzingen naar numerieke methoden (kras voorbeeld; blz.99: "We used a program based on the method of Powel"; is er maar één en is het niet M.J.D. Powell?). Hoewel zelf niet werkzaam op het terrein van normalisatie, kan ik mij voorstellen dat het boek voor diegenen die er mee werken duidelijke praktische adviezen bevat. Veel wiskundige problemen zijn nog blijven liggen (vgl. b.v. blz.182); ook de auteur stelt vast dat dit onderwerp zich nog leent voor vruchtbaar verder onderzoek. Samenvattend: een helder geschreven monografie over een belangwekkend onderwerp, die de waarde en mogelijkheden van toegepast beslis kundig onderzoek duidelijk illustreert.

J. Kriens, K.H.T.

A.P.J. Abrahamse, G.J. van Driel & C. van Ravenswaaij

STATISTIEK, Een kennismaking met achtergronden en methoden.

Van Gorcum & Comp. bv, Assen, 1983, 228p, ISBN 90-232-2007-2, f 37,50

Alvorens meer algemeen op de verdiensten en gebreken van het boek in te gaan, geef ik er de voorkeur aan eerst over elk hoofdstuk afzonderlijk iets op te merken.

Na een inleiding met een aantal historische accenten worden in hoofdstuk 2 enige basisbegrippen behandeld waarbij de behandeling van verzamelingen en gebeurtenissen op deze plaats toch wat vreemd aandoet. Indeling bij het hoofdstuk over kansrekening ligt meer voor de hand.

In hoofdstuk 3 over frekwentieverdelingen en hun grafische voorstellingen komt de samenhang tussen een frekwentiepolygoon en een histogram niet goed naar voren. Als gevolg hiervan is niet duidelijk dat ook bij frekwentiepolygonen met frekwentiedichtheden moet worden gewerkt.

Hoofdstuk 4 over maatstaven voor speciale aspecten van verdelingen bevat een behandeling van lokatiemaatstaven en variatie- en spreidingsmaatstaven. Scheefheidsmaatstaven komen niet aan de orde. De formules en symbolen hebben betrekking op populaties ofschoon wordt uitgegaan van een groep waarnemingsuitkomsten en wordt geschreven dat in plaats van μ men ook wel het symbool $\bar{x}$ tegenkomt dat meestal voor een steekproef wordt gebruikt. Bij de formule voor de standaarddeviatie wordt hierdoor gedeeld door n en niet door $n-1$. Ook in latere hoofdstukken over de stochastiek blijft dit zo, zodat geen gebruik kan worden gemaakt

van de t-verdeling.

Hoofdstuk 5 over associatie is totaal afwijkend ten opzichte van vergelijkbare boeken en ook zeer origineel. Bij deze behandeling van de samenhang tussen nominale eigenschappen wordt het onderscheid tussen marginale samenhang en direkte (partiële) samenhang duidelijk geïllustreerd.

Hoofdstuk 6 geeft een zeer beknopte behandeling van de bekende indexcijfers van Laspeyres, Paasche en Fisher, waarbij opvalt dat in de gebruikte symbolen geen plaats is ingeruimd voor het aangeven van basis en beschouwde periode, hetgeen tot verwarring aanleiding kan geven.

Alhoewel eerder het harmonisch gemiddelde is behandeld, wordt het bij de behandeling van de indexcijfers van Paasche niet gebruikt. Dat aandacht wordt besteed aan het sociaal-ekonomisch jargon is uitstekend maar toch behoren ekonomen meer over belangrijke ekonomische indexcijfers te weten dan in het boek is opgenomen. Hoofdstuk 7 is weer zeer origineel en geeft een specifieke behandeling van begrippen uit de bevolkingsstatistiek, gegeneraliseerd tot kollekties met in- en uittrekking.

In de inleiding in korrelatie en regressie van hoofdstuk 8 doen de meetkundige intermezzo's vreemd aan. Ook de behandeling van multikollineariteit is in tegenstelling tot de opzet van het boek erg technisch zonder goede illustraties met voorbeelden.

Hoofdstuk 9 handelt over dekompositie van tijdreeksen. De vraag is of in de tijd van ARIMA-modellen de behandeling niet zou kunnen worden beperkt tot het laten zien hoe in principe een reeks voor een seizoen kan worden gekorrigeerd.

In hoofdstuk 10 over kansrekening wordt een aardige synthese tot stand gebracht tussen de subjektivistische en de frekwentistische kansdefinitie d.m.v. het gebruik van het eenheidsvierkant.

In hoofdstuk 11 over kansvariabelen en kansverdelingen behandelt geen verdelingsfuncties en kansfuncties in het algemeen, maar begint abrupt met een summiere behandeling van de uniforme, de alternatieve en de binominale verdeling. De normale verdeling wordt alleen behandeld als benadering van de binominale verdeling, waardoor hij als zelfstandige continue theoretische verdeling niet goed uit de verf komt. Op blz. 173 worden de begrippen schatter en rake schatting gebruikt, zonder dat expliciet wordt gezegd wat er mee wordt bedoeld.

Hoofdstuk 12 bevat een uiteenzetting over het hoe en waarom van steekproeven, over steekproefverdelingen (in het boek trekkingsverdelingen genoemd) en over de centrale - limietstelling.

In hoofdstuk 13 over toetsen van hypothesen wordt aan de hand van leuke voorbeelden het verschil tussen linkseenzijdig, rechtseenzijdig en tweezijdig toetsen aangegeven. Naast toetsen over een gemiddelde en een proportie komen toetsen over regressiecoëfficiënten aan de orde.

In hoofdstuk 14 wordt ingegaan op het schatten van een populatiegemiddelde, een populatieproportie en een regressielijn.

Evenals in hoofdstuk 13 wordt uitsluitend gebruik gemaakt van de normale verdeling en niet van de t-verdeling, zodat of σ bekend wordt verondersteld of alleen gebruik kan worden gemaakt van grote steekproeven.

Dit is een grote beperking daar elk standaard computerprogramma voor regressie gebruik maakt van de t-verdeling.

In bovenstaande hoofdstuksgewijze behandeling komen zowel positieve als negatieve punten naar voren. Hetzelfde blijft gelden als men het boek in zijn geheel probeert te beoordelen. Men kan duidelijk merken dat meerdere personen aan de totstandkoming van het boek hebben bijgedragen vanwege het heterogene karakter. Sommige stukken zijn tamelijk ingewikkeld en technisch, andere daarentegen helemaal niet. In het voorwoord wordt vermeld dat het accent meer op de ontwikkeling van inzicht dan op het in detail uiteenzetten van statistische technieken is gelegd. Naar mijn gevoel is men hier zeker in geslaagd in zoverre hiermee wordt bedoeld het door middel van goede voorbeelden intuïtief duidelijk maken hoe bepaalde methoden kunnen worden gebruikt. Anderzijds is bijvoorbeeld bij het regressiemodel noodzakelijk duidelijk te maken onder welke veronderstellingen dit kan worden gebruikt. Dit gebeurt onvoldoende. De verspreide behandeling van de regressierekening in de hoofdstukken 8, 13 en 14 lijkt niet bevorderlijk voor

een goed begrip van dit voor ekonomen zo belangrijke onderwerp. Het boek bevat met zijn 212 pagina's waarschijnlijk ook teveel verschillende onderwerpen zodat een aantal maar zeer oppervlakkig kan worden behandeld. Een cursus voor ekonomen behoeft toch wat meer diepgang met dan maar wat minder onderwerpen.

Al met al bracht het boek na lezing bij mij tegenstrijdige gevoelens teweeg. Het scoort hoog op originaliteit van sommige hoofdstukken en de gebruikte voorbeelden maar anderzijds zijn er een aantal zwakke punten zoals boven weergegeven. Overigens moet worden bedacht dat de gehanteerde hoofdstuksgewijze behandeling de negatieve punten wellicht enigszins benadrukt. Het is in ieder geval een lezenswaardig boek voor iedereen die wil leren werken met statistische gegevens met name op het terrein van de ekonomie.

F.T.M. Klijn, UvA.

R.E. Kirk

Experimental Design: Procedures for the Behavioral Sciences. Second Edition.

Brooks/Cole Publishing Company, Monterey, 1982, 911p, ISBN 0-8185-0286-X, \$57.50

De eerste editie uit 1968, behoorde in het begin van de zeventiger jaren, samen met Winer (1971) tot een van de meest aangehaalde boeken op het gebied van de variantieanalyse in de psychologische onderzoeksliteratuur. Ondanks het feit dat sommige hoofdstukken uit de eerste druk nagenoeg onveranderd zijn overgenomen, kan men stellen dat deze tweede druk een aanzienlijke verandering heeft ondergaan. Opvallend is dat dit boek in een overzichtelijke druk en met een duidelijke notatie is uitgevoerd.

Kirk heeft gepoogd de benadering van variantieanalyse via het algemene lineaire model te integreren met de meer traditionele uitleg van variantieanalyse, zoals o.a. in de eerste druk en in Winer (1971) te vinden is.

In de hoofdstukken 1 en 2 wordt een overzicht gegeven van de basisbegrippen van experimentele proefopzetten en variantieanalyse. In hoofdstuk 3 wordt net zoals in de eerste druk op duidelijke wijze het verschil tussen apriori en aposteriori toetsen uitgelegd en komen achtereenvolgens de verschillende multipale toetsprocedures aan bod. Persoonlijk vind ik het echter jammer, dat in dit hoofdstuk de procedure van Scheffé niet expliciet gepresenteerd wordt als onderdeel van de simultane toetsprocedure van Gabriel. Overigens komt de simultane toetsprocedure wel beknopt aan de orde in hoofdstuk 8 bij de bespreking van de interpretatie van simpele en interactie effecten.

Na een nagenoeg onveranderde behandeling van het model met een onafhankelijke variabele in hoofdstuk 4, wordt in hoofdstuk 5 het algemene lineaire model gepresenteerd. Op duidelijke wijze legt Kirk uit hoe het variantieanalyse model door middel van het algemene lineaire model geformuleerd kan worden. Deze formulering leidt tot de overeenkomst van de variantieanalyse en regressieanalyse modellen. Aan de hand van overzichtelijke schema's geeft Kirk aan hoe het variantieanalyse model te schrijven is als een regressieanalyse model met "dummy" variabelen. Tenslotte wordt het probleem van de oplosbaarheid van de vergelijkingen aan het eind van hoofdstuk 5 uitgebreid behandeld. Deze uitleg kan niet zonder enige kennis van matrix operaties worden begrepen. Een beknopte appendix over matrix algebra is daarom door Kirk toegevoegd. Bij de bespreking van andere variantieanalytische modellen wordt telkens een paragraaf besteed aan de algemene lineaire model benadering.

In de hoofdstukken 6, 9 en 11 worden uitgebreid de herhaalde metingen modellen behandeld, samen met de bijbehorende problemen, zoals de noodzakelijke en voldoende voorwaarden waaraan de populatiecovariantie matrix moet voldoen om een exacte F-toets te kunnen uitvoeren. De rest van het boek wordt besteed aan gekruiste factoriële modellen, Latin-square modellen, modellen met "confounded" interactie-effecten en covariantieanalyse.

Het boek omvat een uitgebreid scala van onderwerpen. Als tekstboek voor een cursus variantieanalyse in een opleiding in een der sociale wetenschappen lijkt het boek dan ook zeer geschikt. In tegenstelling tot Winer (1971) die zich in enkele

paragrafen laat verleiden tot het geven van formules die zonder een diepgaande kennis van de statistiek niet te begrijpen zijn, is de moeilijkheidsgraad in het gehele boek van Kirk ongeveer even hoog. Bij elk hoofdstuk zijn oefenopgaven toegevoegd.

Toch zouden enkele kanttekeningen geplaatst kunnen worden. Bij een uitgebreide behandeling van aspecten van variantieanalyse zou ook een bespreking van variantieanalyse op dichotome en rangorde gegevens horen. Immers, gegevens uit experimenten voldoen vaak niet aan de gestelde assumpties. In de eerste druk kwam wel een beknopte bespreking van niet-parametrische toetsen voor. Mede omdat Kirk toch matrix algebra hanteert, zou een behandeling van multivariate variantieanalyse voor de hand liggen, en tenslotte zou in deze tijd een Bayesiaanse benadering van variantieanalyse in een dergelijk boek niet misstaan. Hoewel zo'n Bayesiaanse benadering nogal gecompliceerd is, zou een eenvoudige uitleg, zoals bijvoorbeeld gegeven wordt door Lewis (1982) goed mogelijk zijn.

Tenslotte is de prijs nogal aan de hoge kant en zou een goedkope studenten uitgave deze barrière om het boek als tekst voor een cursus te gebruiken opheffen.

Referenties.

LEWIS, Ch., Bayesian methods for analysis of variance. In: G. Keren (ed.) *Statistical and Methodological Issues in Psychology and Social Sciences Research*. New Jersey, Laurence Erlbaum Ass. 1982.

WINER, B.J., *Statistical Principles in Experimental Design*. New York, McGraw-Hill, 1971.

M.P.F. Berger, KHT.

Jacqueline Meulman

Homogeneity Analysis of incomplete data

DSWO Press, 1982, ISBN 90-6695-001-3, f23,-

Het onderhavige werk is een gecombineerde onderneming om enerzijds het-voor niet wiskundigen vaak moeilijke-werk dat door de vakgroep datatheorie van de R.U.Leiden gedaan is op het gebied van de homogeniteitsanalyse op de wijze van een inleidend leerboek te presenteren, en anderzijds verslag te doen van een exploratief onderzoek naar het effect van ontbrekende gegevens op de conclusies van een homogeniteitsanalyse.

Het werk valt uiteen in drie delen: (i) een kort literatuuroverzicht over de behandeling van ontbrekende gegevens in de multivariate analyse; (ii) een theoretische uiteenzetting over homogeniteitsanalyse in één en in meerdere dimensies; (iii) de behandeling van onvolledige data, op zijn beurt uiteenvallend in een theoretisch en een toegepast gedeelte, dat zich in hoofdzaak bezighoudt met computersimulaties, aangevuld met een analyse van een empirische dataset.

De inleiding, hoewel goed gestructureerd, staat toch een beetje los van het eigen werk. Het onderzoek dat direct relevant is voor het eigen onderzoek wordt pas in het slothoofdstuk behandeld, terwijl de uitvoerige aandacht voor ML-schaters bij ontbrekende gegevens verder in het boek eigenlijk niet meer relevant is. Een tweede vb: er wordt gesuggereerd dat, wanneer ontbrekende data niet geschat worden, ofwel de strategie van "pairwise deletion" ofwel de strategie van "listwise deletion" wordt toegepast (p_{10}), terwijl bijv. bij de berekening van Benzécri-afstanden, waar in het verdere werk (te)veel aandacht aan wordt besteed, geen van beide strategieën wordt toegepast: zodra er enige informatie over een onderzoeksobject beschikbaar is, is de Benzécri-afstand tot alle andere objecten goed gedefinieerd.

In het theoretische gedeelte wordt de centrale idee van de homogeniteitsanalyse uiteengezet: hoe kan de informatie die in een aantal variabelen aanwezig is met minimaal verlies in een afgeleide variabele worden samengevat. Door het toelaten van niet-lineaire transformaties van de variabelen is de techniek bij uitstek geschikt voor de analyse van kwalitatieve variabelen: de resultaten van de homogeniteitsanalyse zijn invariant onder willekeurige 1-1 transformaties van de variabelen.

De auteur getroost zich veel moeite om - in het spoor van Gifi (1981) - duidelijk te maken dat door de uitgevoerde transformaties - de categorie-kwantificaties - enerzijds de discriminatie tussen de geobserveerde objecten gemaximaliseerd wordt, en anderzijds de homogeniteit van de getransformeerde variabelen (d.i. de gemiddelde covariantie) eveneens gemaximaliseerd wordt. De oplossing van beide problemen blijkt identiek te zijn: de geïnduceerde scores van de objecten en de categorie-kwantificaties zijn proportioneel met de eerste (d.i. belangrijkste) linker en rechter singuliere vector van een matrix H die rechtstreeks uit de gegevens geconstrueerd kan worden. De maat van homogeniteit is het kwadraat van de eerste singuliere waarde van H . Afhankelijk van de gekozen proportionaliteitsconstante - de normalisatie van de eindconfiguratie - krijgt men twee aantrekkelijke interpretatiemogelijkheden: ofwel is de kwantificatie van een categorie het centroid van de geïnduceerde scores van alle objecten die de betreffende categorie 'gekozen' hebben, ofwel - als duale interpretatie - is de geïnduceerde score van een object het centroid van de categoriekwantificaties waardoor dit object wordt gekarakteriseerd.

De meerdimensionale homogeniteitsanalyse (hfst.3) blijkt een relatief eenvoudige veralgemening te zijn van de unidimensionale analyse: de (vector-)object-scores en categoriekwantificaties zijn proportioneel met de p belangrijkste singuliere vectoren van dezelfde matrix H , waardoor de aantrekkelijke eigenschap ontstaat dat oplossingen in een lage dimensionaliteit genest zijn in oplossingen van hogere dimensionaliteit.

Het is verdienenstelijk van de auteur dat zij duidelijke verbanden legt tussen de homogeniteitsanalyse enerzijds, en technieken als niet lineaire principale componenten analyse en de ontvouwingstechniek (Heiser, 1981) anderzijds. De hoofdstukken 2 en 3 zijn, op een paar punten na, van zodanige kwaliteit dat ze als leerboek voor de homogeniteitsanalyse kunnen doorgaan. Toch een paar opmerkingen: (i) het dubbel centreren van de similariteitsmatrix P_0 (p.22) wordt uitgevoerd met het - juiste - argument dat daardoor de triviale oplossing van constante scores en categoriekwantificaties wordt vermeden. Er wordt echter niet aangetoond - noch expliciet, noch door verwijzing - dat door deze operaties de andere karakteristieken van de matrix onaangetast blijven. (ii) Hoewel wordt aangetoond dat in de unidimensionale homogeniteitsanalyse de index van homogeniteit proportioneel is met de grootste eigenwaarde van de correlatiematrix van de getransformeerde variabelen, wordt het theoretisch belangrijke resultaat dat een analoog resultaat - bijvoorbeeld de som van de p grootste eigenwaarden van een correlatiematrix - niet geldig is bij een p -dimensionale analyse zeer onduidelijk behandeld.

Even terzijde kan nog opgemerkt worden dat het lezen van deze allesbehalve eenvoudige theorie onnodig bemoeilijkt wordt door tikfouten in formules en tabellen. Bovendien - maar dit kan niet op de rekening van de auteur geschreven worden - vermoed ik dat de gebruikte programmatuur nogal te lijden heeft van foutieve afrondingen: in de eerste regel van tabel 2.4 (p.28) is de gewogen som der categoriekwantificaties niet nul - zoals vereist - maar -0.042. Dit betekent dat, als voor een klein probleem - 10 objecten, 5 variabelen en 12 categorieën - het effect van de afronding al in het tweede beteknend cijfer te merken is, resultaten van analyses bij omvangrijke problemen met de nodige omzichtigheid moeten behandeld worden. Het zou aanbeveling verdienen de HOMALS-programma's toch eens nauwkeurig op hun numerieke stabiliteit te bekijken.

In het derde deel - hoofdstukken 4 tot 10 - wordt het eigenlijke onderwerp van het boek behandeld: homogeniteitsanalyse bij ontbrekende data. In de theoretische analyse wordt op elegante manier aangetoond dat de schattingsprocedure een tamelijk eenvoudige veralgemening is van de procedure bij volledige data: in het geval van schatting van de objectscores worden de elementen van de al vermelde similariteitenmatrix gewogen met $(m_i m_j)^{-1/2}$ waarbij m_i het aantal variabelen is waarop object i is geobserveerd, ditgeen betekent dat de effectieve observaties relatief zwaarder worden gewogen. Dit is eigenlijk de enige manier waarop in de homogeniteitsanalyse ontbrekende data echt uit de analyse geweerd worden, en waar ze dus geen invloed hebben op het eindresultaat. Deze benadering wordt de passieve benadering genoemd. Daarnaast onderscheidt de auteur twee 'actieve' benaderingen, waarbij eigenlijk informatie wordt toegevoegd: in de 'enkelvoudige' benadering

wordt kunstmatig per variabele een categorie ('missing') toegevoegd, met de implicatie dat het zinvol is objecten als gelijkend te beschouwen omdat ze beide niet zijn ondergebracht c.q. onder te brengen zijn in het klassificatieschema dat door de overige categorieën is gedefinieerd. In de "meervoudige" strategie wordt voor elk ontbrekend gegeven een nieuwe categorie gecreëerd. Door de centroid afbeeldingen zullen deze categoriekwantificaties steeds samenvallen met de score van het geassocieerde object, en bijgevolg niets bijdragen aan de verliesfunctie. De index van homogeniteit zal er echter in vergelijking met de passieve en enkelvoudige strategie door geflatteerd worden. Deze meervoudige kwantificatie zorgt er ook voor dat objecten met ontbrekende observaties steeds minder gaan lijken op welk ander object dan ook, en in het licht van de ontvouwingstechniek die ook hier van toepassing is, kan verwacht worden dat zulke objecten en hun geassocieerde categorie in de onderliggende multidimensionale ruimte extreem zullen afgebeeld worden. Door de afbeelding in een ruimte van lage dimensionaliteit (bv. 1) kunnen deze extreme posities zodanig gaan domineren dat ze bepalend zijn voor de unidimensionale oplossing, zodat een niet-informatief artefact ontstaat. Deze conclusies liggen een beetje voor de hand, en de auteur kondigt ze ook aan (P.53), waarbij de cruciale vraag natuurlijk is onder welke omstandigheden dit artefact zal opduiken. Op p.53 wordt in het algemeen beweerd dat dit afhangt van de structuur van de data, zonder enige verdere aanzet tot theoretisering. In hoofdstuk 8 wordt een klein maar interessant theoretisch resultaat geboekt: indien een object aanwezig is met ontbrekende - dus unieke - categorieën op alle variabelen zal de tegenstelling tussen dit object en alle andere objecten de unidimensionale analyse totaal bepalen. In het afsluitende hoofdstuk krijgen we dan te lezen dat dit artefact onder de vorm van uitbijters zal optreden indien er minstens een object is waarvoor de proportie ontbrekende observaties groter is dan de homogeniteitsindex voor de hele dataset, waarbij het argument waarop deze conclusie berust ontbreekt.

De analyses die door de auteur zijn uitgevoerd hebben betrekking op zg. gauges, d.w.z. datastructuren waarvan de eigenschappen gekend zijn. Een gebruikte gauge is afkomstig van Coombs (1964) en betreft de antwoorden van individuen op een aantal 5-categoriale variabelen (items) met geordende categorieën, waarbij het onderliggende model het beste kan beschreven worden met een unidimensionaal latentcontinuum waarbij de itemkarakteristieke functie eentoppig is, of in deterministische termen: itemcategorieën en personen kunnen afgebeeld worden als punten op een continuum, en de persoon kiest voor elk item die categorie waarvan de afbeelding het dichtste ligt bij zijn eigen afbeelding (ideaalpunt). Indien de data voorgesteld worden in een binaire personen x categorieën matrix, dan bestaat er een unieke permutatie van de rijen, zodanig dat in elke kolom de 1-en aansluiten. Dit is de zgn. Petrie-eigenschap (Heiser, 1981). Zo'n matrix is een gauge (maatstaf) voor het onderliggend model. Bovendien, en dit is belangrijk voor de data-analist, zal de homogeniteitsanalyse deze permutatie vinden en zullen de categoriekwantificaties monotone transformaties zijn van de categorierangnummers. Als derde criterium - naast de Petrie-eigenschap en de monotoniteit der transformaties - om te besluiten of een empirische datastructuur aan de eigenschappen van een gauge voldoet, moeten de categoriekwantificaties in de tweede dimensie een kwadratische transformatie zijn van de eerste dimensie. Bewijs van deze stelling geeft ze niet, terwijl Heiser (1981) alleen het bewijs geeft voor binaire variabelen, en dat in het speciale geval dat alle toegelaten antwoordpatronen met dezelfde frekwentie voorkomen.

Deze drie criteria - waarbij de kwadratische transformatie-eis nogal losjes vertaald wordt in 'hoe fijze vorm' van de categoriekwantificaties en objectscores in het vlak van de eerste en tweede dimensie - worden onderzocht voor de voornoemde gauge ongeval van volledige en van onvolledige data (7%, 14% en 20% op toevalsbasis ontbrekende data). Twee soortgelijke gauges, doch hier met een flinke hoeveelheid ruis in de data worden eveneens onderzocht. Over de bespreking die de auteur aan de analyses wijdt het volgende:

(i) Er is een zeer groot onevenwicht tussen de hoeveelheid tabellen en figuren die gereproduceerd worden en het belang dat ze hebben: meer dan 60 pagina's worden besteed aan figuren en tabellen, waarvan vele niet eens belangwekkend

genoeg zijn om in een appendix opgenomen te worden.

(ii) De resultaten op zichzelf zijn niet erg belangwekkend, en misschien is dit niet verwonderlijk, want er werd nauwelijks met hypothesen gewerkt. Het enige dat van belang kan zijn - doch een verdere en meer gedetailleerde exploratie zou wenselijk zijn - is dat in de deterministische gauge de meervoudige kwantificatie strategie zelfs met 20% ontbrekende data nog goed presteert, terwijl het in de probabilistische gauges misgaat bij 16% en 14%. Een meer gedetailleerd Monte-Carlo onderzoek en een theoretische exploratie zouden hier op hun plaats geweest zijn.

(iii) De stabiliteit van de categoriekwantificaties wordt met de bootstrap-methode onderzocht. Het elimineren van de meest extreme bootstrap resultaten (hoeveel?) om 80% betrouwbaarheidsintervallen af te bakenen, verdient m.i. toch wat meer toelichting dan de twee regels die er aan gewijd worden (p.67). De verschillen tussen de categoriekwantificaties worden stabiel genoemd indien hun 80% betrouwbaarheidsintervallen niet overlappen. Deze werkwijze zou te rechtvaardigen zijn indien de categoriekwantificaties onafhankelijk van elkaar waren, doch dat zijn ze niet. Men zou gemakkelijk een eerste indruk van de covariantie tussen de bootstrap schatters kunnen krijgen door voor alle paren categorieën te tellen hoe vaak de bootstrap schatters de voorspelde monotoniciteit schenden. De onduidelijkheid van de bijbehorende plaatjes met over elkaar heen getekende cijfers laat dit niet toe.

In het 9e hoofdstuk wordt een empirische gauge geanalyseerd: de Münsingen-Rain data waarin de variabelen de aan- of afwezigheid zijn van 70 culturele voorwerpen in 59 graven en waarbij geprobeerd wordt uitsluitend met behulp van deze informatie de graven te seriëren. In modeltermen kunnen deze data verklaard worden met een unidimensionaal latent continuum dat monotoon is met de tijd, en waarbij afwezigheid van een voorwerp indicatief is voor een graf dat dateert van vóór dan wel van na de tijd dat het deponeren van een dergelijk voorwerp in zwang was. De aanwezigheid van een voorwerp in twee graven daarentegen wijst op de temporele nabijheid van die twee graven. Aanwezigheid is dus een homogene categorie, terwijl afwezigheid dubbelzinnig is, en dat is de reden om afwezigheid als ontbrekende observatie op te vatten. Deze data vormen een gauge omdat op grond van andere archeologische evidentie tot een aanvaardbare seriatie kon worden besloten. De resultaten van de analyse zijn fascinerend: de enkelvoudige strategie presteert slecht zoals voorspeld, en dit resultaat is een overtuigende illustratie van Coombs' (1964) principe dat data de neerslag zijn van een theorie, en geeft Heiser (1981, hfst.4) gelijk waar hij het automatische coderingsprincipe van 'dédoublement' aanvecht. Even fascinerend zijn de goede resultaten van de passieve zowel als van de meervoudige strategie. Het contrast met de degeneratie in twee gevallen met de probabilistische gauge en de meervoudige strategie is groot. Er wordt even gewezen op een mogelijke verklaring - het al dan niet scalaïr zijn van een frekwentievector - doch een meer gedetailleerde exploratie van dit probleem zou de data-analist in het veld een nuttig handvat kunnen bieden bij de keuze van zijn analyse-instrument.

Met één opmerking van de auteur ben ik het helemaal oneens: op p.162 wordt gesteld (hoewel niet met grote stelligheid) dat ook zonder de onafhankelijke seriatie op grond van andere archeologische gegevens en theorieën de uitgevoerde analyses een aanvaardbare seriatie opleveren. Doch in dat geval is niet alleen de vraag aan de orde hoe men de beste seriatie kan vinden, doch ook de theoretisch belangrijke vraag of graven op grond van hun inhoud aan cultuurvoorwerpen geseerieerd kunnen worden. De (voorlopig?) nogal losjes gehanteerde criteria van de hoofdzervorm en het ongeveer Petrie-zijn van de gepermuteerde datamatrix zijn niet erg overtuigend zolang niet is aangetoond dat ze niet haalbaar zijn onder belangrijke concurrerende modellen, bijvoorbeeld een model dat totale onafhankelijkheid tussen tijdsperiode en cultuurinhoud van het graf postuleert. Men bedenke dat Guttman bij de scalogramanalyse oorspronkelijk een Rep-coëfficiënt van .85 voorstelde om schalen van niet-schalen te scheiden, terwijl bij toevalsgegevens een verwachte Rep-coëfficiënt van meer dan .90 geen zeldzaamheid is (Torgerson 1958, hfst.12). Een kwantificatie van de mate waarin de eigenschappen van de gauge in empirische datastructuren benaderd worden en een statistische

evaluatie van die kwantificatie ontbreken vooralsnog bij de homogeniteitsanalyse.

Referenties.

COOMBS, C.H., *A Theory of Data*. New-York, Wiley, 1964.

GIFI, A., *Nonlinear Multivariate Analysis*. Leiden, Dept. of Data Theory, 1981.

HEISER, W. *Unfolding Analysis of Proximity Data*. Dissertatie. Leiden, 1981.

TORGERSON, W.S., *Theory and Methods of Scaling*. New-York, Wiley, 1958.

N. Verhelst, R.U.U.

W.E. Biles & J.J. Swain

Optimization and Industrial Experimentation

John Wiley & Sons, New York, 1980, 368p, \$ 70,45

Het boek bevat 7 hoofdstukken met de volgende titels.

1. Introduction
2. Fundamental Statistical Concepts
3. Experimental Design Fundamentals
4. Fundamentals of Optimization
5. Optimization via Experimentation
6. Optimization and Experimentation with Physical Processes
7. Optimization and Computer Simulation Experiments

De hoofdstukken 1 t/m 4 behandelen bekende onderwerpen aangaande optimaliserings-technieken met restricties en proefopzetten. Pas in hoofdstuk 5 komen methoden aan bod waarbij de auteurs op zeer verdienstelijke wijze laten zien hoe men optimaliseren en experimenteren kan combineren, rekening houdend met experimentele fouten.

Het door de auteurs gestelde probleem kan bijvoorbeeld als volgt aan de hand van 2 onafhankelijke variabelen X_1 en X_2 geformuleerd worden. Het gaat om de primaire variabele Y_0 met de restricties Y_1 en Y_2 .

$$Y_0 = f(X_1, X_2) + \epsilon_0 = \beta_0 + \beta_1 X_1^2 + \beta_2 X_2^2 + \epsilon_0 \quad (0, 0.25)$$

$$Y_1 = g(X_1, X_2) + \epsilon_1 = \gamma_0 + \gamma_1 X_1^2 + \gamma_2 X_2^2 + \epsilon_1 \quad (0, 0.25)$$

$$Y_2 = h(X_1, X_2) + \epsilon_2 = \delta_0 + \delta_1 X_1^2 + \delta_2 X_2^2 + \epsilon_2 \quad (0, 0.25)$$

Het optimaliseringsprobleem is als volgt:

$$\text{Maximaliseer } f(X_1, X_2) = E(Y_0)$$

met de restricties

$$g(X_1, X_2) = E(Y_1) \leq 25$$

$$h(X_1, X_2) = E(Y_2) \leq 27$$

$$X_1 \geq 0, \quad X_2 \geq 0$$

De term $\epsilon_j \quad (0, 0.25)$ duidt op een procesfout met een gemiddelde van 0 en een variatie van 0.25. Het gaat er nu om een set (X_1, X_2) waarden te vinden die aan bovenstaande eisen voldoen.

De experimentele en optimalisatie procedures worden dan op een zodanige algemene wijze beschreven dat ze bij veel verschillende probleemformuleringen kunnen worden toegepast.

In hoofdstuk 6 van het boek worden vervolgens een aantal praktische voorbeelden uitvoerig beschreven.

Samenvattend vormt het boek een waardevolle aanvulling op de reeds beschikbare optimalisatie en statistische literatuur.

P.M. Upperman, Philips.

Symposium Statistische Software

Symposiumbundel, Utrecht, 1983, 232p, ca. f25,-

Deze bundel met 16 bijdragen is uitgegeven ter gelegenheid van het Symposium Statistische Software 1983, dat op 2 september 1983 in Utrecht werd gehouden. Zoals reeds is aangekondigd in het VVS bulletin 11 (november 1983) is de bundel herdrukt en kunnen ook belangstellenden die het symposium niet hebben bijgewoond van de inhoud kennis nemen door een exemplaar te bestellen bij het Centrum voor Data-Analyse te Utrecht. De uitgave is keurig verzorgd en komt op het eerste gezicht bij mij over als een luxe editie van een themanummer van Kwantitatieve Methoden. Dat daarbij het VVS-embleem ontbreekt en is vervangen door de fraaie SSS83 logo spreekt vanzelf, maar dat daarnaast de redacteur(en) niet op het omslag en de titelpagina zijn vermeld is toch wel jammer. De lijst van bijdragen, die in bovengenoemd bulletin was opgenomen, zullen we hier omwille van de ruimte niet herhalen. Alle auteurs zijn (of waren tot voor kort) werkzaam aan een Nederlandse universiteit; ongeveer tweederde is afkomstig uit Amsterdam of Groningen. De vier bijdragen die een door de auteurs ontwikkeld softwarepakket behandelen zijn afkomstig uit diverse plaatsen, maar niet uit Groningen of Amsterdam. De Groningers schrijven vooral in de categorie "Beleid ten aanzien van software". Het zal allemaal wel toevallig zijn.

Niet alle artikelen zijn even hoog van kwaliteit, zoals te verwachten is bij een dergelijke bundel, maar ieder die belang stelt in statistische software zal een aantal goede, interessante artikelen aantreffen. Historische artikelen met een blik naar de toekomst, zoals die van Molenaar en Stokman, zijn voor bijna iedereen lezenswaard. Voor degenen die zelf software ontwikkelen zijn meerdere bijdragen interessant. Omdat over de human interface naar de software, zoals Molenaar al opmerkt, nog zo weinig is gepubliceerd, is elke alinea in een artikel die hierop betrekking heeft van belang. Zo heb ik al een paar maal het artikel van Verbeek nageslagen, omdat hij een aantal punten die ik zelf ook wel in het hoofd had systematisch op papier heeft gezet (al is het dan met een low-cost matrixprinter).

Waarom evaluatie van statistische programmatuur zo veel werk is wordt duidelijk uit een bijdrage van Brouwer en Debets, die al eerder over dit onderwerp in KM publiceerden.

Er is een aantal bijdragen opgenomen over onderwijs in statistical computing. Ook dat is van belang, zelfs voor het bedrijfsleven.

Al met al geeft de bundel een waardevolle bijdrage aan de literatuur over statistische software in Nederland.

Floor van Nes, CPB.

K. van Harn

Classifying Infinitely Divisible Distributions

M.C. tract 103, Mathematical Centre Amsterdam, 1978.

This monograph starts from two recurrence relations:

$$(1) \quad (n+1)p_{n+1} = \sum_{k=0}^n \alpha_k p_{n-k},$$

$$(2) \quad p_{n+1} = \sum_{k=0}^n \beta_k p_{n-k}.$$

For $\alpha_k \geq 0$ the first relation defines the infinitely divisible (inf div), i.e. the compound Poisson distributions on $\mathbb{Z}_+$, the second, for $\beta_k > 0$, the subclass of compound geometric distributions (renewal sequences).

After a number of plausible but unsuccessful attempts, the author finds constants $c_n(\alpha)$ such that the solutions of

$$(3) \quad c_{n+1}^{(\alpha)} p_{n+1} = \sum_{k=0}^n \gamma_k p_{n-k} \quad (0 < \gamma < 1)$$

for $\gamma_k \geq 0$ interpolate continuously between (1) and (2). All solutions of (3) are shown to be inf div.

A continuous analogue to (3) is found, and it turns out that the analogue to (2) is closely related to Kingman's p -functions and the potential kernels of Berg and Forst.

In connection with decomposition properties of lattice distributions, analogues are proposed for the concepts of stability and self-decomposability (so far only known for absolutely continuous distributions).

The first chapter provides an introduction to inv div distributions, and a review of canonical measures.

The book is well written (maybe slightly overdetailed) and contains 52 references.

F.W. Steutel, T.H.E.

Th.J.v. Eyck, J.J. Groot, P. v. Kampen

Statistiek voor de bedrijfsadministratie

Wolters-Noordhoff, 19 , 454p, ISBN 900131340X

Het boek is volgens de inleiding geschreven voor degenen die iets verder willen komen in de administratie dan alles tot op de cent kloppend hebben. Het is toegespitst op het bedrijfsbeheer en meer in het bijzonder gelden de exameneisen voor het Staatspraktijkdiploma voor Bedrijfsadministratie als leidraad.

Daar er eerder leerboeken voor Statistiek-SPD zijn verschenen, zou men zich kunnen afvragen of dit boek persé nodig is. Het SPD-programma is in de loop der jaren, naast de samenvoeging van Statistiek in het eerste deel, eniger mate uitgebreid. Bovendien zijn de in de literatuur en praktijk gehanteerde technieken en begrippen de afgelopen 10 à 15 jaar op een ander niveau (hoger?) gebracht.

Wie 30 jaar geleden wiskunde heeft gestudeerd op de middelbare school, kan het leerboekje van de eerste klas Mavo nu niet of met moeite volgen. Met de Statistiek gaat het dezelfde weg en het is derhalve goed een poging tot aanpassing te wagen.

Wie als opleider dit boek doorneemt komt al ras tot de konklusie, dat de auteurs er zeer zorgvuldig naar hebben gestreefd het genoemde examenprogramma volledig door te werken - zelfs de "onnauwkeurige getallen" worden met 29 pagina's wel zeer nauwkeurig behandeld. Daarenboven vinden wij ook zeer nuttige stof die men niet in het examenprogramma maar wel in het dagelijks werken tegenkomt, zoals de current en de quick ratio.

Overigens mis ik in hoofdstuk 3 "Verhoudingsgetallen" onder het hoofdje "kenggetallen van de verkoop" het begrip omzetsnelheid. Het is mogelijk dat de bij inkoop toegelichte omloopsnelheid van de voorraad hiervoor mede moet doorgaan, maar door het verschil met de meer gebruikelijke terminologie van "omzetsnelheid" kan de leerling op het verkeerde been komen te staan. De opleider behoeft hem of haar alsdan daarvoor.

Het is overigens juist gezien, dat statistische technieken het best worden duidelijk gemaakt aan de hand van exacte voorbeelden. De leerlingen zijn nu eenmaal als regel debet- en creditgerichte mensen, die - een enkele uitgezonderd - graag werken met veel getallen en weinig woorden.

In hoofdstuk 4 "Indexcijfers" worden Laspeyres en Paasche dan ook zeer getalsmatig duidelijk uit de doeken gedaan. De nadelen van Paasche worden wellicht iets te soepel aangeduid. Dat bij over de gehele linie stijgende prijzen de prijsindex van Paasche kan dalen, mag best breed worden uitgemeten.

Voorts lezen wij dat het basisverleggen bij samengesteld gewogen indices niet anders mogelijk is dan door het wegingsschema aan te passen. Volkomen juist en

wij hopen dat maar, dat toekomstige examencommissies dit ter harte nemen. In op-gaven wordt nogal eens om lapmiddelen gevraagd, terwijl ook de praktijk in deze vaak de nodige inventiviteit vergt. De opleider wijze hierop en kan hierin al zijn creativiteit nog kwijt.

Het hoofdstuk "Frequentieverdelingen" is altijd bijzonder belangrijk, omdat bij de behandeling hiervan de basis wordt gelegd voor het goede begrip van onderwerpen, die hier logischerwijs op volgen zoals de centrummaten, spreidingsmaatstaven, normale verdeling, kansregels en steekproeven. Hoewel naar veler oordeel wellicht wat te uitgebreid wordt ingegaan op het bepalen van het aantal klassen - een gevolg van het overboord gooien van de bekende vuistregel? - wordt in dit hoofdstuk een goede ondergrond gegeven voor het doorwerken van de volgende genoemde hoofdstukken. De overgang naar het volgende chapter "centrummaten" (rekenkundig gemiddelde, modus, mediaan) gaat dan ook als vanzelf. De student met weinig wiskundige ervaring zal wellicht wat schrikken van deze begrippen in symbolen, maar met enige moeite en consistente assistentie van de opleider moet dit eenvoudig te overwinnen zijn. Men is nogal eens geneigd om formules uit het hoofd te leren, terwijl de opleiding erop gericht dient te zijn, dat de begrippen inhoudelijk tot geestelijk eigendom worden gemaakt en dan verkort in symbolen kunnen worden weergegeven.

Deze algemeen stichtende opmerking geldt eens te meer voor het onderwerp spreidingsmaatstaven en in weer mindere mate voor "de normale verdeling". De bekende kromme van Gauss krijgt hierbij de eer van een apart hoofdstuk. Op zich geen luxe, want veel cursisten sukkelen jaren met de toepassing van de standaarddeviatie. Bovendien wordt hier en passant het woord "kans" gebruikt.

Zo belanden we bijna automatisch bij de wiskundige kansregels en steekproeven. Het is een verademing dat hieraan separaat en vrij uitgebreid aandacht wordt geschonken. Deze onderwerpen worden in de meest gebruikte handboeken vaak onderbelicht en verdienen nu juist op grond van hun importantie en vaak voor de student als moeilijk te vatten leerstof extra aandacht.

Bijzonder uitvoerig wordt "Correlatie" uit de doeken gedaan en biedt in feite meer dan benodigd is voor het maken van de examenopgaven. Zo is bijvoorbeeld de afleiding van de methode van de kleinste kwadraten en de correlatiecoëfficiënt naar mijn inschatting niet zo zeer voor het gemiddelde doch eerder voor de echte liefhebber bedoeld.

Hier staat dan weer tegenover, dat de "Tijdreeksanalyse" eenvoudig en vlot leesbaar is. Vooral voor degenen die de kleine lettertjes overslaan en dat zullen er met instemming van de schrijvers velen zijn.

De onderwerpen "Bedrijfseconomische Statistiek" en "Algemeen Economische Statistiek" zijn in opzet beperkt tot de meest gangbare examenstof. Hoewel dit niet is euvel te duiden, kan men zich wel afvragen of in een tijd waarin zoveel gereguleerd wordt aan de hand van statistische gegevens, enige uitbreiding of liever accentverschuiving van de exameneisen naar het praktisch hanteren van de statistiek aanbevelingswaardig zou kunnen zijn.

"Marktonderzoek" staat traditioneel achter in het boek, doch niet als laatste hoofdstuk.

Als laatste onderwerp wordt in hoofdstuk 16 "Statistisch onderzoek" behandeld. De onderscheidene soorten grafieken en tabellen worden veelal in de eerste hoofdstukken van de leerboeken behandeld. Waarom dit onderwerp als besluit is gekozen is mij niet duidelijk.

Dit geldt overigens dan alleen de plaats van het slothoofdstuk, want resumerend mag worden gesteld, dat het hier besproken boek op grond van volledigheid, leesbaarheid en duidelijkheid een hoog cijfer verdient.

Naar de mening van ondergetekende mogen opleiders en vooral ook de studerende de heren Van Eijck, Groot en Van Kampen complimenteren met hun geslaagde poging het statistiekonderwijs te verrijken door het geheel hanteerbaar en begrijpelijk te presenteren.

Peter Sprent

Quick statistics, an introduction to non-parametric methods

Penguin Books, 1981, 264p, £3,95

Dit aardige pocketboekje geeft op elementair niveau een inleiding in de verdelingsvrije methoden. Het is ingedeeld in twee delen.

Het eerste deel (85p) behandelt de benodigde basisbegrippen: enige combinatoriek, kansen, de normale verdeling als benadering voor discrete verdelingen.

Het tweede deel (165p) behandelt alle klassieke rangtoetsen, $n \times m$ tabellen, Kolmogorov-Smirnov e.d.. Hierbij worden vele voorbeelden gegeven. De toon is lichtvoetig. Alleen de elementaire bewijzen worden gegeven. Voor moeilijke bewijzen wordt de lezer verwezen naar de literatuur. Naast toetsen wordt ook veel aandacht gegeven aan de constructie van betrouwbaarheidsintervallen voor plaatsparameters.

In het laatste deel worden de opgaven besproken en summiere tabellen voor de besproken toetsingsgrootheden gegeven.

Een nadeel van het boekje is dat de titels van de hoofdstukken nogal vaag zijn en dat niet eenvoudig is na te gaan, welke toetsen waar behandeld worden. Het is geschreven als leerboek en niet erg geschikt als naslagwerk. Het is dan ook niet duidelijk voor welke groep gebruikers dit boek nuttig is. Misschien kan het nuttig zijn als tekst naast een college over elementaire statistiek. Als tekst om een college op te baseren is het echter minder geschikt omdat de opzet nogal los is en steeds uitgaat van de voorbeelden.

Resumerend, een charmant boekje dat voor een lage prijs een leuke introductie geeft in parametervrije methoden, en dat naast een cursus elementaire statistiek bruikbaar kan zijn.

J.C. van Houwelingen, R.U.U.

G. McLeod

Box Jenkins in Practice

Uitgegeven in eigen beheer door Gwilym Jenkins & Partners Ltd, Lancaster, Engeland, 1982, 430p, £ 20,-

Het beschouwde boek is grotendeels een weergave van de cursus 'Box Jenkins Forecasting; Advanced Theory and Practical Applications' welke door het instituut van Jenkins wordt verzorgd. McLeod heeft het boek geschreven voor mensen die vanuit de praktijk werkzaam zijn en die reeds enige ervaring hebben opgedaan met het bouwen en vooral interpreteren van Box Jenkins modellen ten behoeve van een voorspelproces.

De auteur is er in geslaagd de opbouw van het boek, gezien zijn doelstelling om voor mensen uit de praktijk te schrijven, modulair te houden. Het boek is verdeeld in 14 secties. De laatste 4 secties zijn het omvangrijkst, totaal 350 pagina's, en hierin komen achtereenvolgens de volgende onderwerpen aan de orde, te weten:

1. de univariate stochastische modellen;
2. 'single input - single output' transfer-ruis modellen;
3. interventie modellen;
4. 'multiple input - single output' transfer-ruis modellen.

Deze secties zijn volgens een vast schema opgebouwd. De modelbouwcyclus wordt uitvoerig beschreven en met behulp van schema's en grafieken verduidelijkt. Zo komen stapsgewijs de volgende fasen aan bod: 1. de identificatie van een model; 2. de schatting van een model; 3. de diagnostische verificatie van een model; 4. Het voorspellen met het 'beste' model; en 5. de modelcontrole (monitoring).

Met behulp van het in het boek aanwezige datamateriaal, sectie 6, kan men de diverse verhelderende voorbeelden narekenen. Eén bezwaar tekent zich hierbij aan: men zou eigenlijk moeten beschikken over de programmatuur van Jenkins. Bezit men al dan niet deze programmatuur dan blijken er in de voorbeelden

(helaas) een aantal, soms storende, fouten te zitten.

In dit boek wordt de lezer een groot scala 'handigheidjes' aangereikt. Men treft in dit boek echter geen (statistische) bewijzen aan. Als leerboek en/of theoretisch naslagwerk is het dan ook niet geschikt. Wil men meer inzicht krijgen in de theoretische processen achter de Box Jenkins filosofie dan wordt men doorverwezen naar o.a. het standaardboek van Box en Jenkins (1).

Samenvattend kan gesteld worden dat dit boek voor de praktijk van de voor- spellers met behulp van de Box Jenkins technieken zeer waardevol is.

Referentie.

- (1) BOX, G.E.P. & G.M. JENKINS, *Time Series Analysis, Forecasting and Control*, Holden Day, San Francisco, 1976, (revised edition).

Herman H.M. Visser, PTT.

O. Moeschlin & D. Pallaschke (eds.)

Game theory and Related Topics

North Holland Publishing Company, Amsterdam, New York, Oxford, 1979.

This is the proceedings of a seminar held in Bonn and Hagen in september 1978. The main topic was game theory.

In the first part of the book one can find:

- (i) Contributions in cooperative game theory by Albers, Bruijneel, Weber, Egea, Richter, Neyman and Tauman.
 - (ii) A paper by Armbruster and Böge on Bayesian game theory.
 - (iii) Papers on dynamic games by Groenewegen and Wessels, Rieder, Rothblum and van der Wal.
 - (iv) A paper by Winkels in which an algorithm is described to determine all equilibria of a bimatrix game.
 - (v) A paper by Peleg on game theoretical analysis of voting schemes.
- The second part is devoted to fixed point theory and optimization theory with a contribution of E. Sperner, who held the honorary chairmanship of the seminar and by Fan, Fenske, Filus, Puppe, Robin and Rolewicz.

The third part consists of six papers on measure theoretic concepts and other tools.

The last part deals with mathematical economics with contributions of E. Dierker, H. Dierker and W. Trockel, Eichhorn, Fuchssteiner and Schröder, Greenberg and Müller, Jacobs, Klein, J. Los and M.W. Los, Meister, Postlewaite and Schmeidler.

Also a list of participants is included. Because of the success of this conference, in 1980 a second conference took place which resulted in the proceedings 'Game Theory and Mathematical Economics', O. Moeschlin & D. Pallaschke eds. North Holland, etc., 1981.

S. Tijs, K.U.N.